



ФЕДЕРАЛЬНОЕ ГОСУДАРСТВЕННОЕ КАЗЕННОЕ ОБРАЗОВАТЕЛЬНОЕ
УЧРЕЖДЕНИЕ ВЫСШЕГО ОБРАЗОВАНИЯ
«МОСКОВСКИЙ УНИВЕРСИТЕТ МИНИСТЕРСТВА ВНУТРЕННИХ ДЕЛ
РОССИЙСКОЙ ФЕДЕРАЦИИ ИМЕНИ В.Я. КИКОТЯ»

**В. А. Богаевский,
И. А. Паршутин,
А. Н. Сударик**

**ОБРАБОТКА
ДАННЫХ ПСИХОЛОГИЧЕСКИХ ИССЛЕДОВАНИЙ
С ПОМОЩЬЮ СТАТИСТИЧЕСКИХ ПРОГРАММ**

Учебно-практическое пособие

Москва
2019

ББК 88
Б73

Рецензенты:

*врио заместителя начальника ОМПО УРЛС ГУ
МВД России по Московской области – начальник отделения психологи-
ческой работы, майор внутренней службы **А. В. Пономарев**;
психолог ГРЛС Отдела МВД России по Головинскому району г. Москвы,
старший лейтенант внутренней службы **И. Н. Черницына***

Богаевский В. А.

Б73 *Обработка данных психологических исследований с помощью
статистических программ* : учебно-практическое пособие /
В. А. Богаевский, И. А. Паршутин, А. Н. Сударик. – М. : Москов-
ский университет МВД России имени В.Я. Кикотя, 2019. – 124 с.
ISBN 978-5-9694-0722-0

Учебно-практическое пособие предназначено для курсантов, слушателей, обучающихся по специальности 37.05.02 – Психология служебной деятельности (специализация – Психологическое обеспечение служебной деятельности сотрудников правоохранительных органов), адъюнктов, обучающихся по специальности 19.00.03 – Психология труда, инженерная психология, эргономика, слушателей факультета переподготовки и повышения квалификации, обучающихся по программам профессиональной переподготовки психологов, преподавателей, читающих психологические дисциплины «Экспериментальная психология», «Математические методы в психологии» и «Методология и методы социально-психологического исследования», образовательных организаций высшего образования МВД России.

Издаваемый материал может быть использован обучающимися при написании курсовых и выпускных квалификационных работ в форме дипломной работы, а также диссертаций для обработки и анализа эмпирических данных психологических исследований, в том числе с использованием компьютерных статистических программ, научных статей, докладов, а также преподавателями психологии для проведения практических занятий по соответствующим специализированным дисциплинам и своей научно-исследовательской работы.

ББК 88

ISBN 978-5-9694-0722-0

© Московский университет МВД России
имени В.Я. Кикотя, 2019
© Богаевский В. А., 2019
© Паршутин И. А., 2019
© Сударик А. Н., 2019

СОДЕРЖАНИЕ

ВВЕДЕНИЕ.....	4
1. Измерение в психологическом исследовании.....	5
2. Общие принципы проверки статистических гипотез.....	12
2.1. Понятие нормального распределения.....	12
2.2. Понятие выборки.....	17
2.3. Проверка статистических гипотез.....	18
2.4. Этапы принятия статистического решения.....	21
2.5. Психологические задачи, решаемые с помощью статистических методов.....	22
3. Статистические критерии различия.....	27
3.1. Анализ нормальности распределения.....	27
3.2. Выбор статистического критерия различия.....	38
3.3. Т-критерий Стьюдента (параметрический критерий для связанных и несвязанных выборок).....	49
4. Корреляционный анализ.....	62
5. Вычисление частотных характеристик.....	71
6. Регрессионный анализ.....	78
7. Факторный анализ.....	88
БИБЛИОГРАФИЧЕСКИЙ СПИСОК.....	97
ПРИЛОЖЕНИЯ.....	98
Приложение 1.....	98
Приложение 2.....	105
Приложение 3.....	107
Приложение 4.....	114

ВВЕДЕНИЕ

Математическая статистика в руках психолога может и должна быть мощным инструментом, позволяющим не только успешно лавировать в море экспериментальных данных, но и, прежде всего, способствовать становлению его объективного мышления.

Настоящее учебное пособие призвано решить следующие задачи:

- 1) дать представление об основных статистических процедурах;
- 2) научить проводить первоначальную статистическую обработку данных экспериментальных исследований;
- 3) научить делать правильные психологические выводы на основе результатов статистического анализа;
- 4) грамотно подготавливать данные для работы со статистическим пакетом программ «MS Office Excel» и «Statistica», правильно понимать результаты их работы и результаты проведенного анализа эмпирических данных.

1. Измерение в психологическом исследовании

В ходе психологического исследования изучаемые характеристики могут получать количественное выражение, например баллы по шкалам теста. Полученные количественные данные эксперимента подвергаются затем статистической обработке [2].

Измерение, проводимое в психологическом исследовании, может быть определено как приписывание чисел изучаемым явлениям, которое осуществляется по определенным правилам [2].

Измеряемый объект сравнивается с некоторым эталоном, в результате чего получает его численное выражение. Закодированная в числовой форме информация позволяет использовать математические методы и выявлять то, что без обращения к числовой интерпретации могло бы остаться скрытым. Кроме того, числовое представление изучаемых явлений позволяет оперировать сложными понятиями в более сокращенной форме. Именно этими обстоятельствами объясняется использование измерений в любой науке, в том числе и психологии.

В целом научно-исследовательскую работу психолога, проводящего измерения, можно представить в следующей последовательности [2]:

1. Исследователь (психолог).
2. Предмет исследования (психические свойства, процессы, функции и т. п.).
3. Испытуемый (группа испытуемых).
4. Эксперимент (измерение).
5. Данные эксперимента (числовые коды).
6. Статистическая обработка данных эксперимента.
7. Результат статистической обработки (числовые коды).
8. Выводы (печатный текст: отчет, диплом, статья и т. п.).
9. Получатель научной информации (руководитель курсовой, дипломной или кандидатской работы, заказчик, читатель статьи и т. п.).

Любой вид измерения предполагает наличие единиц измерения. Единица измерения – это та «измерительная палочка», как говорил С. Стивенс, которая является условным эталоном для осуществления тех или иных измерительных процедур [2]. В естественных науках и технике существуют стандартные единицы измерения, например, градус, метр, ампер и т. д.

Психологические переменные за единичными исключениями не имеют собственных измерительных единиц. Поэтому в большинстве случаев значение психологического признака определяется при помощи специальных измерительных шкал.

Согласно С. Стивенсу, существует четыре типа измерительных шкал (или способов измерения) [2]:

- 1) номинативная (номинальная или шкала наименований);
- 2) порядковая (ординарная или ранговая шкала);
- 3) интервальная (шкала равных интервалов);
- 4) шкала отношений (шкала равных отношений).

Все находящиеся в скобках наименования являются синонимами изначальному понятию.

Процесс присвоения количественных (числовых) значений имеющейся у исследователя информации называется кодированием. Иными словами, кодирование – это такая операция, с помощью которой экспериментальным данным придается форма числового сообщения (кода).

Применение процедуры измерения возможно только четырьмя вышеперечисленными способами. Причем каждая измерительная шкала имеет собственную, отличную от других форму числового представления или кода. Поэтому закодированные признаки изучаемого явления, измеренные по одной из названных шкал, фиксируются в строго определенной числовой системе, определяемой особенностями используемой шкалы.

Измерения, осуществляемые с помощью двух первых шкал, считаются качественными, а осуществляемые с помощью двух последних шкал – количественными. С развитием научных познаний все более возрастающее

значение приобретает количественное описание на основе методов измерения. При этом преследуются две конкретные цели:

1. Повышение и оценка степени точности вывода. Количественные данные позволяют по сравнению с качественными описаниями достичь более высокой степени точности и дают при этом возможность для принятия более обоснованных решений.

2. Формулирование законов. Цель каждой науки – описывать через законы существенные отношения между исследуемыми явлениями. Если эти отношения можно выразить количественно в виде функциональных зависимостей, то прогностические возможности сформулированного таким образом закона природы значительно возрастают.

Номинативная шкала (шкала наименований)

Измерение в номинативной шкале состоит в присваивании какому-либо свойству или признаку определенного обозначения или символа (численного, буквенного и т. п.). По сути дела, процедура измерения сводится к классификации свойств, группировке объектов, к объединению их в классы при условии, что объекты, принадлежащие к одному классу, идентичны (или аналогичны) друг другу в отношении какого-либо признака или свойства, тогда как объекты, различающиеся по этому признаку, попадают в разные классы [2].

Иными словами, при измерениях по этой шкале осуществляется классификация или распределение объектов (например, типы акцентуации характера личности) на непересекающиеся классы, группы. Таких непересекающихся классов может быть несколько. Классический пример измерения по номинативной шкале в психологии – разбиение людей по четырем темпераментам: сангвиник, холерик, флегматик и меланхолик.

Номинальная шкала определяет, что разные свойства или признаки качественно отличаются друг от друга, но не подразумевает каких-либо количественных операций с ними. Так, для признаков, измеренных по этой шкале, нельзя сказать, что какой-то из них больше, а какой-то меньше,

какой-то лучше, а какой-то хуже. Можно лишь утверждать, что признаки, попавшие в разные группы (классы), различны. Последнее и характеризует данную шкалу как качественную.

Приведем еще пример измерения в номинативной шкале [2]. Психолог изучает мотивы увольнения с работы:

- а) не устраивал заработок;
- б) неудобная сменность;
- в) плохие условия труда;
- г) неинтересная работа;
- д) конфликт с начальством и т. д.

Самая простая номинативная шкала называется дихотомической. При измерениях по дихотомической шкале измеряемые признаки можно кодировать двумя символами или цифрами, например 0 и 1, или буквами А и Б, а также любыми двумя отличающимися друг от друга символами. Признак, измеренный по дихотомической шкале, называется альтернативным.

В дихотомической шкале все объекты, признаки или изучаемые свойства разбиваются на два непересекающихся класса, при этом исследователь ставит вопрос о том, «проявился» ли интересующий его признак у испытуемого или нет. Например, в исследовании из 30 испытуемых принимали участие 23 женщины, которым можно поставить цифру 0, и 7 мужчин, кодируемых цифрой 1.

Приведем еще примеры, относящиеся к измерениям по дихотомической шкале:

- а) испытуемый ответил на пункт опросника либо «да», либо «нет»;
- б) кто-то проголосовал «за», кто-то «против»;
- в) человек либо «экстраверт», либо «интроверт» и т. д.

Во всех перечисленных случаях получаются два непересекающихся множества, применительно к которым можно только подсчитать количество индивидов, обладающих тем или иным признаком.

В номинативной шкале можно подсчитать частоту встречаемости признака, т. е. число испытуемых, явлений и т. п., попавших в данный класс (группу) и обладающих данным свойством.

Порядковая (ранговая, ординарная) шкала

Измерение по этой шкале расчленяет всю совокупность измеренных признаков на такие множества, которые связаны между собой отношениями типа «больше – меньше», «выше – ниже», «сильнее – слабее» и т. п. Если в предыдущей шкале было несущественно, в каком порядке располагаются измеренные признаки, то в порядковой (ранговой) шкале все признаки располагаются по рангу – от самого большого (высокого, сильного, умного и т. п.) до самого маленького (низкого, слабого, глупого и т. п.) или наоборот.

Типичный и очень хорошо известный всем пример порядковой шкалы – это школьные оценки: от 5 до 1 балла.

В порядковой (ранговой) шкале должно быть не меньше трех классов (групп): например, ответы на опросник: «да», «не знаю», «нет».

Приведем еще пример измерения в порядковой шкале. Психолог изучает социометрические статусы членов коллектива:

1. «Популярные».
2. «Предпочитаемые».
3. «Пренебрегаемые».
4. «Изолированные».
5. «Отвергаемые».

Шкала интервалов (интервальная шкала)

В шкале интервалов, или интервальной шкале, каждое из возможных значений измеренных величин отстоит от ближайшего на равном расстоянии. Главное понятие этой шкалы – интервал, который можно определить как долю или часть измеряемого свойства между двумя соседними позициями на шкале. Размер интервала – величина фиксированная и постоянная на всех участках шкалы.

При работе с этой шкалой измеряемому свойству или предмету присваивается соответствующее число. Важной особенностью шкалы интервалов является то, что у нее нет естественной точки отсчета (ноль условен и не указывает на отсутствие измеряемого свойства).

Так, в психологии часто используется семантический дифференциал Ч. Осгуда, который является примером измерения по интервальной шкале различных психологических особенностей личности, социальных установок, ценностных ориентаций, субъективно-личностного смысла, различных аспектов самооценки и т. п.:

-3	-2	-1	0	1	2	3
абсолютно			не знаю			абсолютно
не согласен						согласен

Шкала отношений (шкала равных отношений)

Шкалу отношений называют также шкалой равных отношений. Особенностью этой шкалы является наличие твердо фиксированного нуля, который означает полное отсутствие какого-либо свойства или признака.

Шкала отношений, по сути, очень близка интервальной, поскольку если строго фиксировать начало отсчета, то любая интервальная шкала превращается в шкалу отношений.

Именно в шкале отношений производятся точные и сверхточные измерения в таких науках, как физика, медицина, химия и др. Приведем примеры: сила гравитации, частота сердцебиения, скорость реакции. В основном измерение по шкале отношений производится в близких к психологии науках, таких, как психофизика, психофизиология, психогенетика. Это обусловлено тем, что очень сложно найти пример психического явления, которое потенциально могло бы отсутствовать в деятельности человека.

Контрольные вопросы и задания

1. Какие измерительные шкалы (способы измерения) Вы знаете?
2. Какие величины измеряет номинативная шкала? Приведите примеры.
3. Какие величины измеряет порядковая шкала? Приведите примеры.
4. Какие величины измеряет шкала интервалов? Приведите примеры.
5. Какие величины измеряет шкала отношений? Приведите примеры.
6. В чем особенность шкалы семантического дифференциала Ч. Осгуда?
7. Какая последовательность действий психолога, проводящего измерение?

2. Общие принципы проверки статистических гипотез

2.1. Понятие нормального распределения

Нормальное распределение играет большую роль в математической статистике, поскольку многие статистические методы предполагают, что анализируемые с их помощью экспериментальные данные распределены нормально. График нормального распределения имеет вид колоколообразной кривой (рис. 1).

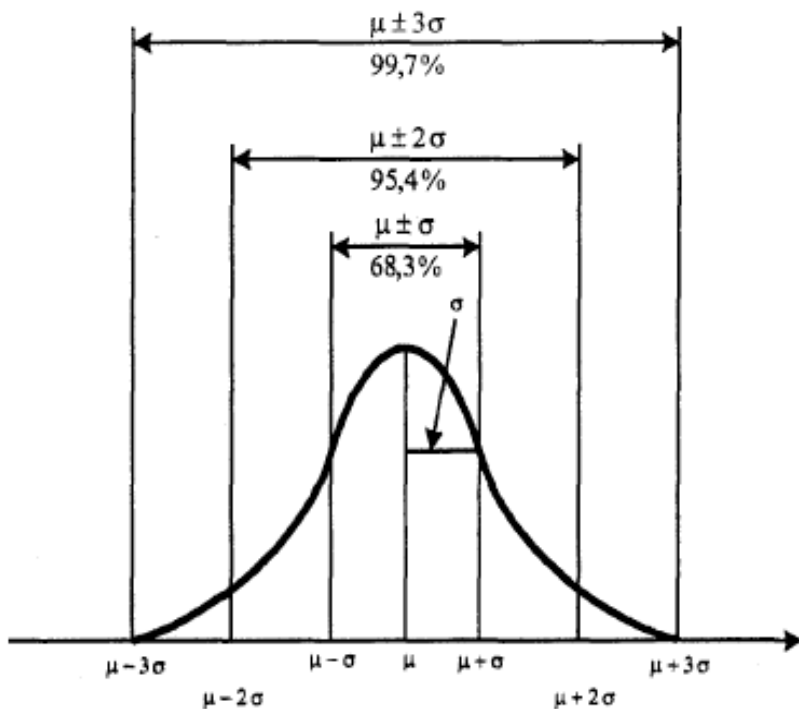


Рис 1. Кривая нормального распределения

Его важной особенностью является то, что форма и положение графика нормального распределения определяются только двумя параметрами: средним значением μ (мю) и стандартным отклонением σ (сигма).

Нормальным такое распределение было названо потому, что оно наиболее часто встречалось в естественнонаучных исследованиях и казалось «нормой» распределения случайных величин.

В реальных психологических экспериментах редко получают данные, распределенные строго по нормальному закону. В большинстве случаев сырые психологические данные часто дают асимметричные, «ненормальные» распределения. По мнению Е. В. Сидоренко, причина этого заключается в самой специфике некоторых психологических признаков, когда от 10 до 20 % испытуемых могут получить либо очень низкие, либо очень высокие оценки по методикам [8]. Распределение таких оценок не может быть нормальным, как бы ни увеличивался объем выборки.

Важнейшими первичными статистиками, по которым проверяется нормальность распределения выборочных эмпирических данных, являются:

а) **средняя арифметическая величина (среднее)** – величина, сумма отрицательных и положительных отклонений выборочных данных от которой равна нулю. В статистике ее обозначают буквой **M** или **X**;

б) **мода** – это числовое значение измеряемой характеристики, которое встречается в выборке наиболее часто. Мода обозначается как **M_o**;

в) **медиана** – это числовое значение измеряемой характеристики, по отношению к которой, по крайней мере, 50 % выборочных значений меньше нее и, по крайней мере, 50 % – больше. Обозначается как **M_e**;

г) **среднее квадратичное отклонение (стандартное отклонение)** – мера разнообразия входящих в группу объектов. Она показывает, на сколько в среднем отклоняется каждая варианта (конкретное значение оцениваемой характеристики) от средней арифметической величины. Обозначается греческой буквой σ (сигма). Чем сильнее разбросаны варианты относительно средней арифметической величины, тем большим оказывается среднее квадратичное отклонение.

Сигма – величина именованная и зависит не только от степени варьирования, но и от единиц измерения. Поэтому по сигме можно сравнивать изменчивость лишь одних и тех же показателей, а сопоставлять сигмы разных признаков по абсолютной величине нельзя. Для того чтобы сравнить по уровню изменчивости признаки любой размерности (выраженные в различных единицах измерения) и избежать влияния масштаба измерения средней арифметической на величину сигмы, применяют коэффициент вариации, который представляет собой, по существу, приведение к одинаковому масштабу величины σ .

д) **коэффициент вариации** – частное от деления сигмы на среднее, умноженное на 100 %. Обозначается коэффициент вариации как CV и вычисляется по формуле:

$$CV = \sigma / M \times 100 \%. \quad (2.1)$$

Для нормального распределения эмпирических данных характерна закономерность, установленная между модой, медианой и средним и выражающаяся в следующем равенстве:

$$M = M_o = M_e. \quad (2.2)$$

Для нормального распределения известны точные количественные зависимости частот и значений, позволяющие прогнозировать появление новых вариантов:

- слева и справа от средней арифметической величины лежит 50 % вариант;
- в интервале от $M - 1\sigma$ до $M + 1\sigma$ лежат 68,3 % всех вариант;
- в интервале от $M - 2\sigma$ до $M + 2\sigma$ лежат 95,4 % вариант;
- в интервале от $M - 3\sigma$ до $M + 3\sigma$ лежат 99,7 % вариант [2].

Таким образом, ориентируясь на эти характеристики нормального распределения, можно оценить степень близости к нему рассматриваемого распределения.

Следующими по важности являются такие первичные статистики, как коэффициент асимметрии и эксцесс.

Коэффициент асимметрии (A) – численная мера скошенности распределения в левую или правую сторону по оси ординат. Коэффициент асимметрии вычисляется по формуле:

$$A = \sum_{i=1}^N f_i (x_i - x)^3 / N \times \sigma^3, \quad (2.3)$$

где f_i – частота x_i -го значения измеряемой характеристики в выборочных данных;

N – объем выборочных данных.

Если правая ветвь кривой распределения длиннее левой (правосторонняя скошенность), говорят о положительной асимметрии, в противоположном случае (левосторонняя скошенность) – об отрицательной. Для нормального (симметричного) распределения коэффициент асимметрии равен нулю ($A=0$).

Эксцесс (E) – количественная мера «горбатости» распределения, показатель островершинности (туповершинности) кривой распределения.

Коэффициент эксцесса вычисляется по формуле:

$$E = \sum_{i=1}^N f_i (x_i - x)^4 / N \times \sigma^4 - 3. \quad (2.4)$$

Кривые, более высокие в своей средней части, островершинные, называются эксцессивными, у них большая величина эксцесса и его значение имеет положительный знак. При уменьшении величины эксцесса кривая становится все более плоской, приобретая вид плато, а затем и седловины – с прогибом в средней части. При этом значение эксцесса

имеет отрицательный знак. Величина эксцесса в нормальном распределении равняется нулю ($E=0$).

Эти параметры позволяют составить первое приближенное представление о характере распределения.

Точную и строгую оценку нормальности распределения можно получить, используя один из существующих методов проверки, например с помощью критерия «хи-квадрат» Пирсона $\chi^2_{\text{эм}}$ [2]. Пример решения такой задачи будет рассмотрен в подразделе 3.1.

Репрезентативность – степень соответствия выборочных показателей генеральным параметрам.

Статистические ошибки репрезентативности показывают, в каких пределах могут отклоняться от параметров генеральной совокупности (от математического ожидания или истинных значений) наши частные распределения, полученные на основании конкретных выборок. Очевидно, что величина ошибки тем больше, чем больше варьирование признака и чем меньше выборка. Это и отражено в формулах для вычисления статистических ошибок, характеризующих варьирование выборочных показателей вокруг их генеральных параметров.

В число первичных статистик входит **статистическая ошибка средней арифметической величины**. Формула для ее вычисления такова:

$$m_M = \pm \sigma / n^{1/2}, \quad (2.5)$$

где m_M – ошибка среднего;

σ – стандартное отклонение;

n – число значений признака.

2.2. Понятие выборки

Выборка – любая группа испытуемых, выделенных из совокупности всех людей, относительно которых ученый намерен сделать выводы при изучении конкретной проблемы (генеральная совокупность) [2].

Выборки называются *независимыми* (несвязанными), если процедура эксперимента и полученные результаты измерения некоторого свойства у испытуемых одной выборки не оказывают влияния на особенности протекания этого же эксперимента и результаты измерения этого же свойства у испытуемых (респондентов) другой выборки.

И, напротив, выборки называется *зависимыми* (связанными), если процедура эксперимента и полученные результаты измерения некоторого свойства, проведенные на одной выборке, оказывают влияние на другую.

Следует подчеркнуть, что одна и та же группа испытуемых, на которой дважды проводилось психологическое обследование (пусть даже разных психологических качеств, признаков, особенностей), все равно оказывается зависимой, или связанной выборкой.

2.3. Проверка статистических гипотез

Полученные в экспериментах выборочные данные всегда ограничены и носят в значительной мере случайный характер. Именно поэтому для анализа таких данных и используется математическая статистика, позволяющая обобщать закономерности, полученные на выборке, и распространять их на всю генеральную совокупность [2].

В начале математической обработки данных экспериментатором выдвигается статистическая гипотеза – формальное предположение о том, что сходство (или различие) некоторых эмпирических данных случайно или, наоборот, неслучайно.

Сущность проверки статистической гипотезы заключается в том, чтобы установить, согласуются ли экспериментальные данные и выдвинутая гипотеза, допустимо ли отнести расхождение между гипотезой и результатом статистического анализа экспериментальных данных за счет случайных причин? Таким образом, статистическая гипотеза – это научная гипотеза, допускающая статистическую проверку, а математическая статистика – это научная дисциплина, задачей которой является научно обоснованная проверка статистических гипотез.

При проверке статистических гипотез используются два понятия: так называемая нулевая (обозначение H_0) и альтернативная гипотеза (обозначение H_1).

Принято считать, что нулевая гипотеза H_0 – это гипотеза о сходстве, а альтернативная H_1 – гипотеза о различии. Таким образом, принятие нулевой гипотезы H_0 свидетельствует об отсутствии различий, а гипотезы H_1 – о наличии различий.

При проверке гипотезы экспериментальные данные могут противоречить гипотезе H_0 , тогда эта гипотеза отклоняется. В противном случае, т. е. если экспериментальные данные согласуются с гипотезой H_0 , она не отклоняется. Часто в таких случаях говорят, что гипотеза H_0 принимается.

Отсюда видно, что статистическая проверка гипотез, основанная на экспериментальных данных, неизбежно связана с риском (вероятностью)

принять ложное решение. При этом возможны ошибки двух родов. Ошибка первого рода произойдет, когда будет принято решение отклонить гипотезу H_0 , хотя в действительности она оказывается верной. Ошибка второго рода произойдет, когда будет принято решение не отклонять гипотезу H_0 , хотя в действительности она будет неверна. Очевидно, что и правильные выводы могут быть приняты также в двух случаях. Вышесказанное представлено в табл. 2.1.

Таблица 2.1

Статистические ошибки

Результат проверки гипотезы H_0	Возможные состояния проверяемой гипотезы	
	<i>Верна гипотеза H_0</i>	<i>Верна гипотеза H_1</i>
<i>Гипотеза H_0 отклоняется</i>	Ошибка первого рода	Правильное решение
<i>Гипотеза H_0 не отклоняется</i>	Правильное решение	Ошибка второго рода

Не исключено, что исследователь может ошибиться в своем статистическом решении. Поскольку исключить ошибки при принятии статистических гипотез невозможно, то необходимо минимизировать возможные последствия, т. е. принятие неверной статистической гипотезы. В большинстве случаев единственный путь минимизации ошибок заключается в увеличении объема.

При обосновании статистического вывода следует решить вопрос, где же проходит линия между принятием и отвержением нулевой гипотезы? В силу наличия в эксперименте случайных влияний эта граница не может быть проведена абсолютно точно. Она базируется на понятии уровня значимости.

Уровнем значимости называется вероятность ошибочного отклонения нулевой гипотезы. Или, иными словами, уровень значимости – это вероятность ошибки первого рода при принятии решения. Для

обозначения этой вероятности, как правило, употребляют латинскую букву p .

Исторически сложилось так, что в прикладных науках, использующих статистику, и в частности в психологии, считается, что низшим уровнем статистической значимости является уровень $p=0,05$; достаточным – уровень $p=0,01$ и высшим – уровень $p=0,001$.

Величины 0,05, 0,01 и 0,001 – это так называемые стандартные уровни статистической значимости. При статистическом анализе экспериментальных данных психолог в зависимости от задач и гипотез исследования должен выбрать необходимый уровень значимости. Как видим, здесь наибольшая величина, или нижняя граница уровня статистической значимости, равняется 0,05 – это означает, что допускается пять ошибок в выборке из ста элементов (испытуемых) или одна ошибка из двадцати элементов. Считается, что ни шесть, ни семь, ни большее количество раз из ста мы ошибиться не можем. Цена таких ошибок будет слишком велика.

2.4. Этапы принятия статистического решения

Принятие статистического решения разбивается на семь этапов [2].

1. Формулировка нулевой и альтернативной гипотез.

2. Определение объема выборки n . Для психологических исследований рекомендуется использовать экспериментальную и контрольную группы так, чтобы численность обоих сравниваемых групп была не менее 30–35 испытуемых в каждой. Планирование эксперимента должно включать в себя учет как объема выборки, так и ряда ее особенностей. Так, в психологических исследованиях важно соблюдение однородности выборки. Оно означает, что психолог, изучая, например подростков, не может включать в эту же выборку взрослых людей. Экспериментальная выборка должна представлять (моделировать) генеральную совокупность, поскольку выводы, полученные в эксперименте, предполагается в дальнейшем перенести на всю генеральную совокупность (репрезентативность).

3. Выбор соответствующего уровня значимости или вероятности отклонения нулевой гипотезы. Это может быть величина меньшая или равная 0,05. В зависимости от важности исследования можно выбрать уровень значимости в 0,01 или даже в 0,001. Как правило, если выборка испытуемых менее 100 человек, используется уровень значимости, равный 0,05.

4. Выбор статистического метода, который зависит от типа решаемой психологической задачи.

5. Вычисление соответствующего эмпирического значения по экспериментальным данным согласно выбранному статистическому методу.

6. Определение статистической достоверности полученных значений, соответствующих выбранному вами уровню значимости ($p=0,05$, $p=0,01$, $p=0,001$).

7. Формулировка принятия решения (выбор соответствующей гипотезы H_0 или H_1).

2.5. Психологические задачи, решаемые с помощью статистических методов

Начать с анализа первичных статистик надо еще и по той причине, что они весьма чувствительны к наличию выпадающих вариантов. На практике же, очень большие эксцессы и асимметрия часто являются индикатором ошибок при подсчетах вручную или ошибок при введении данных через клавиатуру при компьютерной обработке. Грубые промахи при введении данных в обработку можно обнаружить, если сравнить величины сигм у аналогичных параметров. Выделяющаяся величиной сигма может указывать на ошибки.

Если обнаружены «подозрительные» значения, то необходимо принять обоснованное решение об их выбраковке. Его можно принять, используя достаточно мощный параметрический критерий t . Он рассчитывается по следующей формуле:

$$t = (V-M)/\sigma \geq t_{st}, \quad (2.6)$$

где t – критерий выпада;

V – выпадающее значение признака;

M – средняя арифметическая величина признака для всей группы, включающей артефакт;

σ – стандартное отклонение по выборке данных;

t_{st} – стандартные значения критерия выпадов, определяемые для трех уровней доверительной вероятности по табл. 2.1 (см. приложение 1).

Смысл критерия в том, чтобы определить, находится ли данная варианта в интервале, характерном для большинства членов выборки, или же вне его.

Допустим, нами принят уровень значимости 0,05 (доверительная вероятность 0,95), а значение критерия составило 1,5. Поскольку 95 % вариантов лежат в пределах $M \pm 1,96\sigma$ (1,5 меньше 1,96), следовательно, дан-

ная варианта лежит в указанном интервале. Если же значение критерия больше, например 2,4, то это означает, что данное значение не относится к анализируемой совокупности (выборке), включающей 95 % вариантов, а есть проявление иных закономерностей, ошибок и пр. и поэтому должно быть исключено из рассмотрения.

Например, в эксперименте вы предлагаете решать мыслительные задачи и регистрируете в числе других параметров время решения. При просмотре данных обнаруживаете, что у одного из испытуемых время решения заметно больше, чем у остальных. Это бывает связано с тем, что вместо решения очередной задачи испытуемый начинает «искать закономерность более широкого плана», «выводить общий принцип» или нечто подобное. Об этом он может сообщить экспериментатору, но может и не сообщать. Понятно, что время решения конкретной задачи при этом может сильно отличаться от средней величины. В этом случае вы окажетесь перед необходимостью принять обоснованное решение – включать данное значение в дальнейшую обработку или нет.

Предположим, в вашем эксперименте были получены следующие значения некоторого параметра: 10, 20, 20, 30, 30, 40, 40, 50, 210. В нашем примере $n=9$. Вычисляем: $M=50$; $\sigma=61$. Можно ли считать значение 210 выпадающим?

$$t = \frac{210 - 50}{61} = 2.62 ;$$

t_{st} (по табл. 2.1) = 2,43 (для $p>0,05$).

Следовательно, значение 210 может считаться выпадающим и должно быть исключено из дальнейшей обработки.

После исключения выпадающих значений первичные статистические параметры вычисляются заново.

Существует правило, согласно которому все расчеты вручную должны выполняться дважды (особенно ответственные – триж-

ды), причём желательно разными способами, с вариацией последовательности обращения к числовому массиву. Оценка генеральной совокупности на основе выборочных данных недостаточно точна, имеет некоторую большую или меньшую ошибку. Такие ошибки, представляющие собой ошибки обобщения, экстраполяции, связанные с перенесением результатов, полученных при изучении выборки, на всю генеральную совокупность, называются ошибками репрезентативности.

Подавляющее большинство задач, решаемых психологом в эксперименте, предполагает те или иные сопоставления. Это могут быть сопоставления одних и тех же показателей в разных группах испытуемых или, напротив, разных показателей в одной и той же группе.

Для определения степени эффективности каких-либо воздействий (обучение, тренировка, тренинг, инструктаж и т. п.) сравниваются показатели «до» и «после» этих воздействий. Например, сравниваются показатели уровня тревожности у подростков «до» и «после» психотренинга, что позволяет определить его эффективность. Или в лонгитюдном исследовании сопоставляются результаты у одних и тех же испытуемых по одним и тем же методикам, но в разном возрасте, что позволяет выявить временную динамику анализируемых показателей.

Два выборочных распределения сравниваются между собой или с теоретическим законом распределения, чтобы выявить различия или, напротив, сходство в типах распределений. Например, сравнение распределений времени решения простой и сложных задач позволит построить классификацию задач и типологию испытуемых.

В целом психологические задачи, решаемые с помощью методов математической статистики, условно можно разделить на несколько групп:

1. Задачи, требующие установления сходства или различия.
2. Задачи, требующие установления связи между данными.
3. Задачи, требующие определения частоты изучаемого свойства.
4. Задачи, требующие установления прогноза.
5. Задачи, требующие группировки и классификации данных.

Контрольные вопросы и задания

1. Нарисуйте кривую нормального распределения данных. Какие ее характеристики?
2. Назовите первичные статистики, по которым проверяется нормальность распределения выборочных данных.
3. Дайте определение понятию «выборка».
4. Какое отличие между независимыми и зависимыми выборками?
5. Какие статистические гипотезы Вы знаете? Приведите примеры формулировок гипотез.
6. Что такое статистическая ошибка? Какие статистические ошибки встречаются в исследованиях?
7. Назовите этапы принятия статистического решения. В чем особенность каждого из этапов?
8. Какие задачи можно решать с помощью методов математической статистики?
9. Что мы называем «выскакивающей» вариантой? Какой алгоритм их отсеивания из выборочных данных?

3. Статистические критерии различия

3.1. Анализ нормальности распределения

Методы статистического анализа можно подразделить на две группы – параметрические и непараметрические. Важным условием, определяющим возможность применения параметрических методов, является подчинение анализируемых данных закону нормального (Гауссова) распределения, в то время как для непараметрических методов выполнения этого условия не требуется [8].

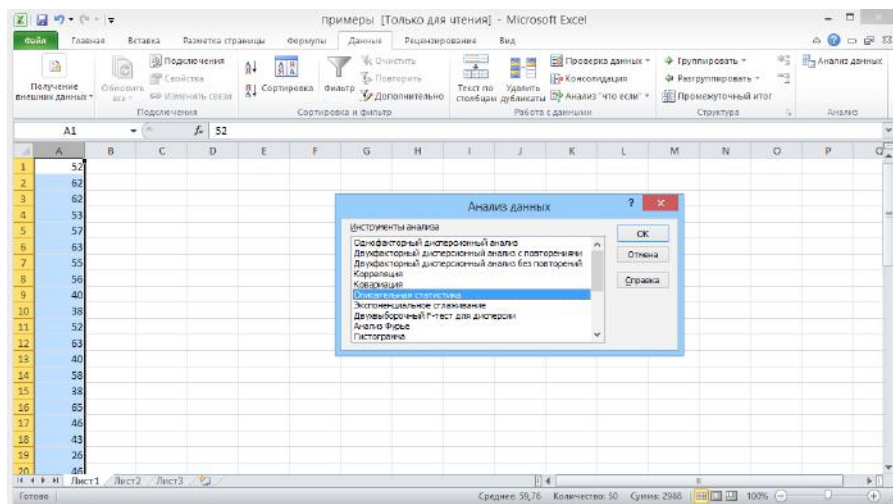
Существует большое количество методов проверки нормальности распределения, но ни один из них не является универсальным [5, 9].

При нормальном распределении, которое симметрично, значения медианы и среднего арифметического будут одинаковы, а значения асимметрии (Skewness) и эксцесса (Kurtosis) равны нулю. Если средняя арифметическая больше медианы, а коэффициент асимметрии > 0 , то распределение имеет правостороннюю асимметрию (скошено вправо). При левосторонней асимметрии средняя арифметическая меньше медианы, а коэффициент асимметрии < 0 . По величине коэффициента эксцесса говорят об островершинном (Kurtosis > 0) или плосковершинном (Kurtosis < 0) распределении. Однако ситуаций, когда средняя арифметическая равна медиане, а коэффициенты асимметрии и эксцесса равны нулю, практически не встречается, поэтому необходимо решить, какие отклонения от идеального сценария допустимы для того, чтобы считать распределение полученных данных нормальным или близким к нормальному [5, 10].

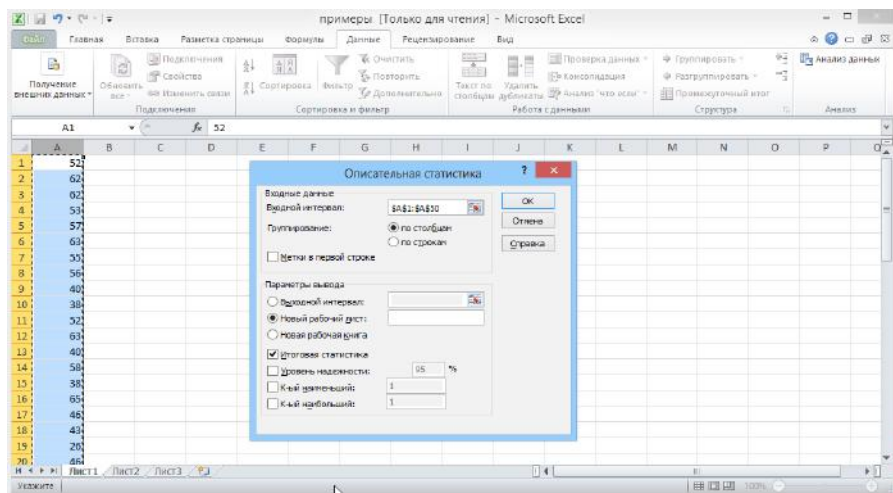
Статистический пакет «MS Office Excel» позволяет провести формальный тест на нормальное распределение на основе вышеупомянутых показателей асимметрии (скошенности) и эксцесса.

Допустим, исследователь должен выяснить, укладываются ли значения шкалы «Самооценки умственных способностей» методики Дембо-Рубинштейн в кривую нормального распределения на примере 50 случаев (количество испытуемых).

Для этого необходимо на панели инструментов выбрать «Данные» > «Анализ данных» > «Описательная статистика»:



В открывшемся диалоговом окне выставите галочку напротив показателя «Итоговая статистика» и введите промежуток массива данных в строке «Входного интервала»:



Итоговый результат статистического анализа будет выглядеть следующим образом:

The screenshot shows a Microsoft Excel window titled 'примеры [Только для чтения] - Microsoft Excel'. The 'Данные' (Data) tab is active. A table of statistical results is displayed in the worksheet, with columns labeled 'A' and 'B'.

Строка	А	Б
1	Среднее	59,76
2	Стандартная ошибка	2,123178583
3	Медиана	60
4	Мода	70
5	Стандартное отклонение	15,02728773
6	Дисперсия выборки	225,8187753
7	Экссесс	0,446738006
8	Асимметричность	0,10879051
9	Интервал	74
10	Минимум	20
11	Максимум	100
12	Сумма	2968
13	Счет	56

Обратите внимание, что итоговый расчет будет автоматически перенесен на новый лист Excel (в нашем случае Лист 5). В приведенных результатах видно, что: среднее (средняя арифметическая) и медиана значимо не отличаются. Также при нормальном распределении данных исследования показатели асимметрии и эксцесса по модулю не должны превышать трех ошибок асимметрии и трех ошибок эксцесса, соответственно. Соблюдение этих условий позволяет сделать вывод о нормальном характере распределения данных исследования.

При помощи формулы ошибки асимметрии (3.1) и ошибки эксцесса (3.2) необходимо провести соответствующие вычисления:

$$Sa = 6(n-1) / (n+1)(n+3), \quad (3.1)$$

$$Se = 24n(n-2)(n-3) / (n+1)^2(n+3)(n+5). \quad (3.2)$$

После применения формул итоговые значения нужно сравнить с данными эксцесса и асимметричности, уже полученными при помощи Excel:

Асимметричность (As)	0,11
Ошибка асимметрии (Sa)	0,34
Эксцесс (Ex)	0,45
Ошибка эксцесса (Se)	0,66

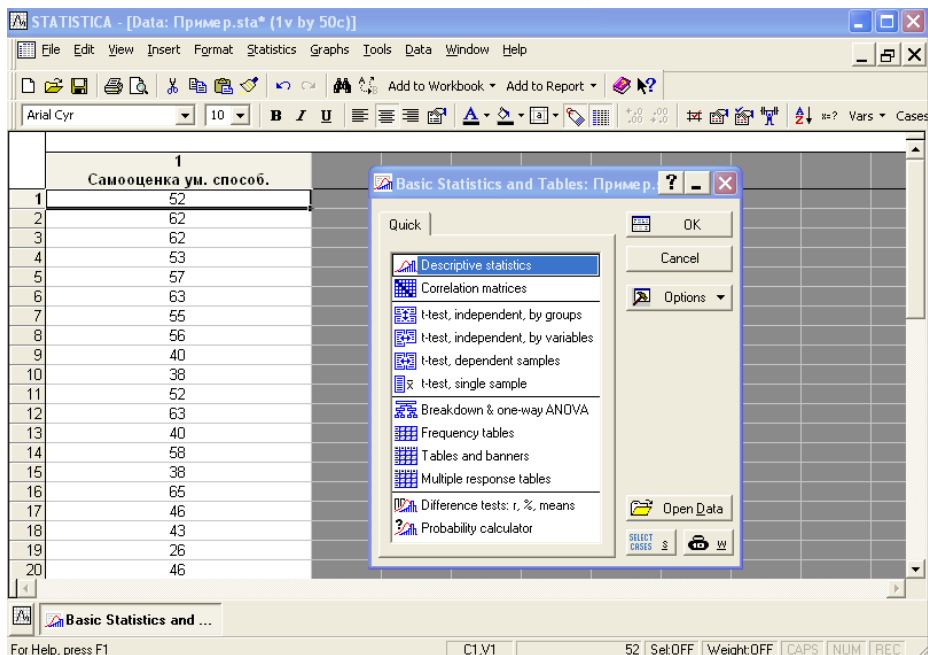
Таким образом, полученные результаты позволяют нам сделать вывод о нормальном характере распределения данных методики изучения самооценки.

При помощи статистической программы «Statistica» также возможно определить нормальность распределения данных исследования, но уже несколькими способами.

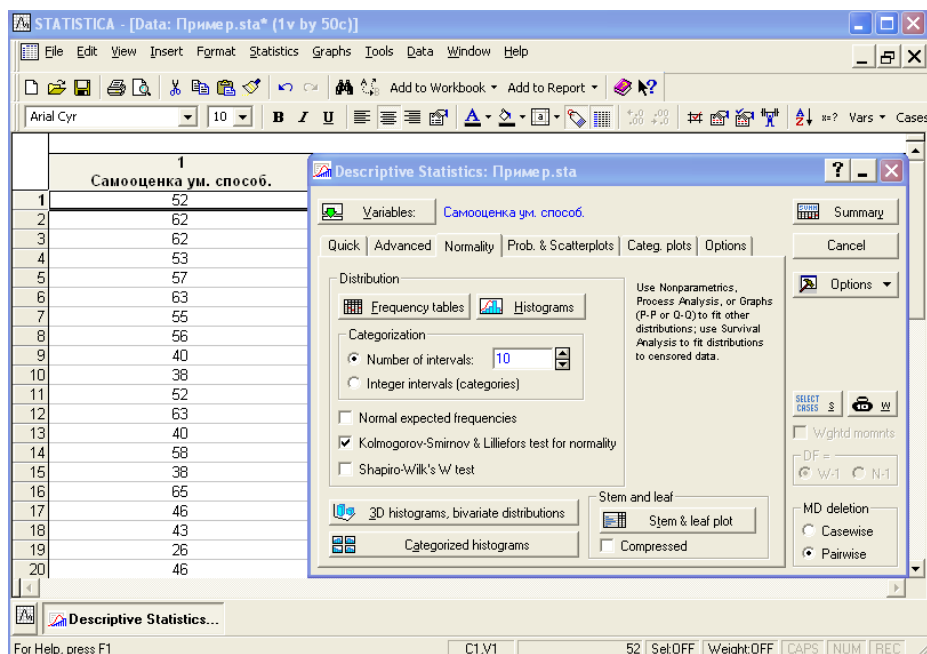
Первый способ

При проверке нормальности распределения кривая нормального распределения может быть наложена на гистограмму. В качестве примера возьмём предыдущую задачу со значениями шкалы «Самооценки умственных способностей» методики Дембо-Рубинштейн).

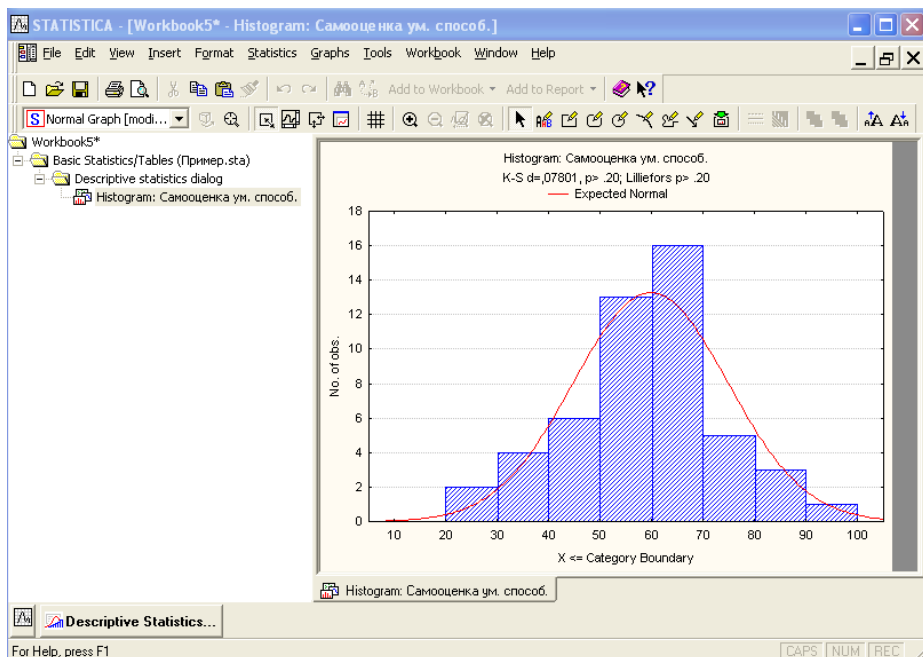
После проведения подготовительных работ с документом (см. приложение 4) выберите «Statistics» («Статистика») > «Basic statistics/Tables» («Основные статистики и таблицы») > «Descriptive statistics» («Описательная статистика»):



В открывшемся диалоговом окне нажмите на «Descriptive statistics», после чего в модуле «Normality» («Нормальное распределение») выставите галочку напротив статистического критерия различий Колмогорова-Смирнова и Лиллифорса. В закладке «Variables» («Переменные») выберите изучаемый параметр (в нашем примере значения «Самооценки умственных способностей») и нажмите на «Histograms» («Построение гистограмм»):



При построении графиков на нормальность распределения программа Statistica 6.0 рассчитывает значения коэффициентов Колмогорова-Смирнова и Лиллифорса. Эти критерии основаны на нулевой гипотезе о том, что данная выборка получена из генеральной совокупности, имеющей нормальное распределение. Если уровень значимости больше 0,05 (т. е. превышает 5 %), то можно принять нулевую гипотезу – или, строго говоря, нет оснований ее отвергнуть [9].

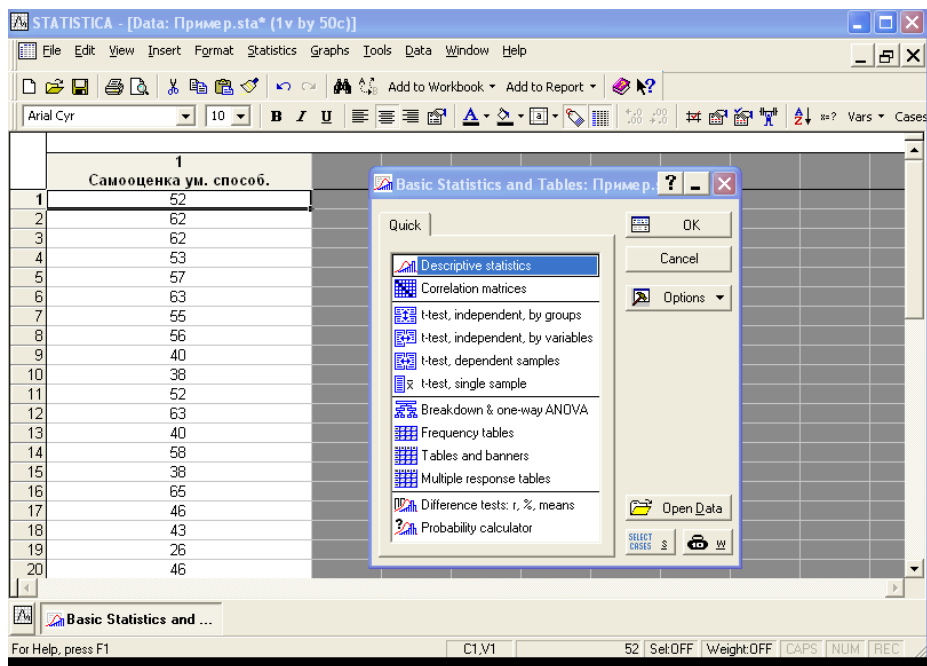


В нашем примере уровень значимости показателей коэффициента Колмогорова-Смирнова и Лиллифорса оказался больше 0,2 ($p > 0,20$), что свидетельствует о нормальном распределении полученных данных.

Другим статистическим критерием, используемым в Statistica 6.0 для проверки нормальности распределения данных, является тест на проверку распределения с помощью критерия Шапиро-Уилка (Shapiro-Wilk). Данный тест предпочтителен для использования при небольших объемах выборки, а с увеличением количества наблюдений достоверность его снижается.

Второй способ

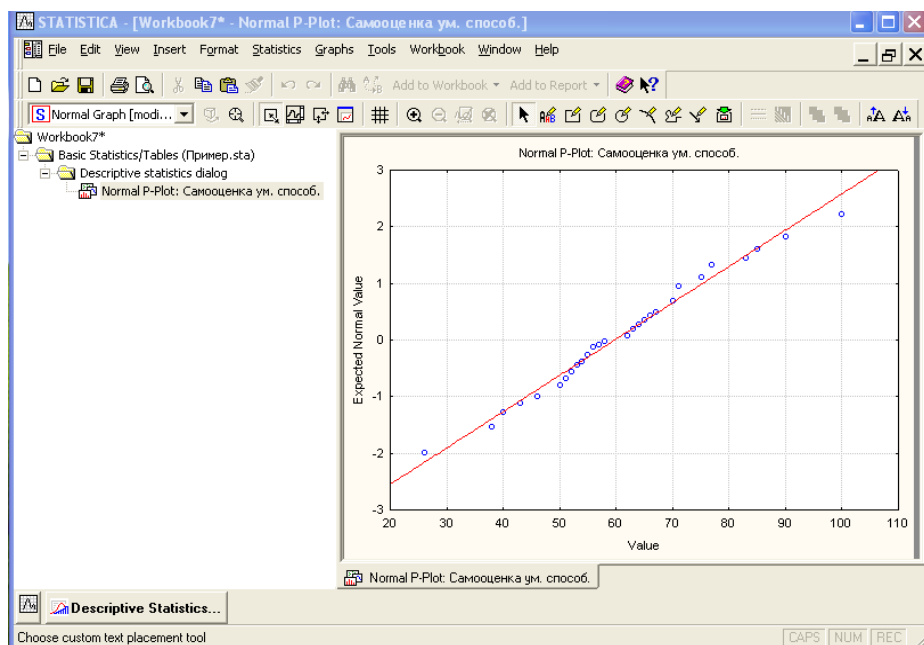
Одним из методов исследования нормальности распределения является также построение графиков. На графике даются координаты фактических значений переменных и теоретические значения, вычисленные при условии нормальности распределения (линия). Такой график изображает зависимость ожидаемых нормальных частот значений признака от их реальных частот. Очевидно, что если между наблюдаемым и ожидаемым распределениями нет никакой разницы, точки на этом графике выстроятся строго вдоль прямой. Если иначе, то они образуют фигуру, отличную от прямой [10].



Для построения графика такого типа необходимо в модуле «Descriptive Statistics» перейти на закладку «Prob. & Scatterplots» («Вероятностные

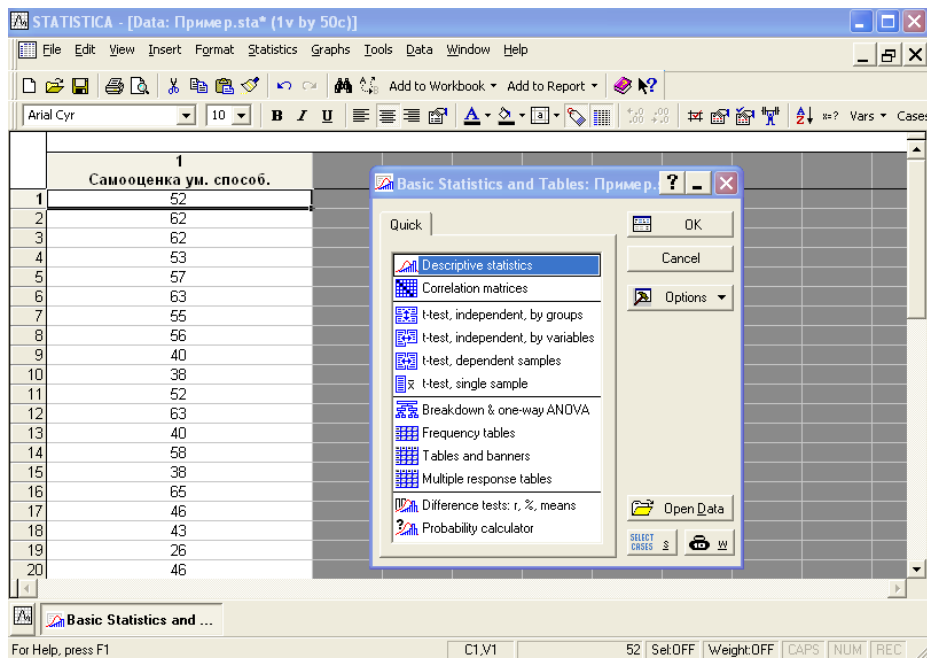
графики и диаграммы рассеяния») и нажать на кнопку «Normal probability plot» («График нормальных вероятностей»).

В результате появится график, точки на котором плотно выстраиваются вдоль теоретически ожидаемой прямой, что еще раз подтверждает вывод о нормальности распределения данных шкалы «Самооценки умственных способностей» методики Дембо-Рубинштейн.



Третий способ

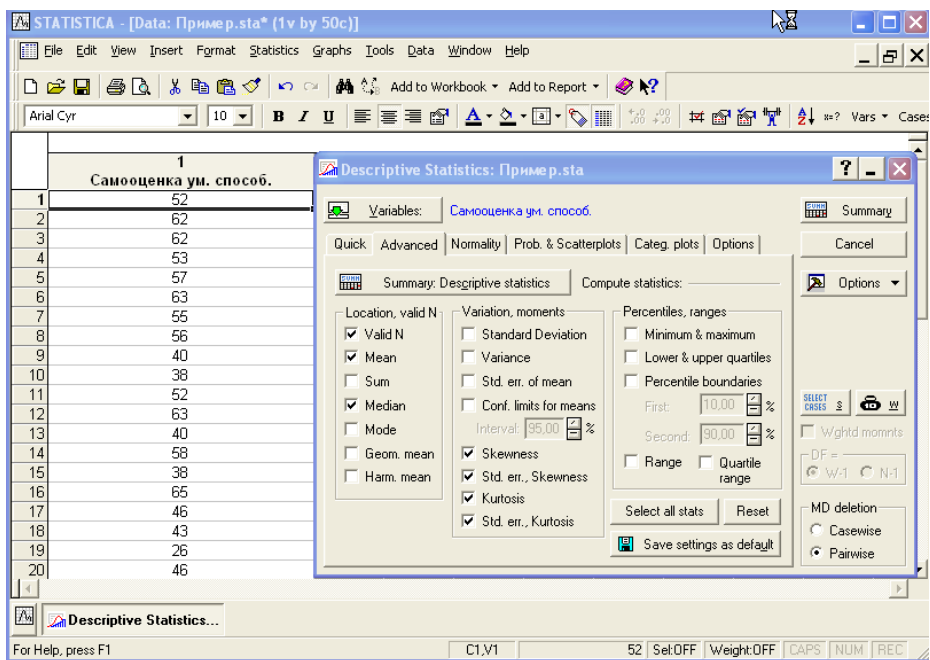
Statistica 6.0 позволяет провести формальный тест на нормальное распределение на основе вышеупомянутых показателей асимметрии (скошенности) и эксцесса. Для этого необходимо выбрать «Statistics» («Статистика») > «Basic statistics/Tables» («Основные статистики и таблицы») > «Descriptive statistics» («Описательная статистика»):



В открывшемся диалоговом окне нажмите на «Descriptive statistics», после чего в закладке «Advanced» («Углубленные методы анализа данных») выставите галочки напротив следующих статистических показателей:

- Valid N – объем выборки;
- Mean – арифметическая средняя;
- Median – медиана;
- Skewness – коэффициент асимметрии;
- Std. err., Skewness – стандартная ошибка коэффициента асимметрии;
- Kurtosis – коэффициент эксцесса;
- Std. err., Kurtosis – стандартная ошибка коэффициента эксцесса.

Также необходимо в закладке «Variables» («Переменные») выбрать учаемый параметр (в нашем примере значения «Самооценки умственных способностей»), после чего нажмите на «Summary».



Итоговый результат статистического анализа будет выглядеть следующим образом:

Variable	Valid N	Mean	Median	Skewness	Std. Err. Skewness	Kurtosis	Std. Err. Kurtosis
Самооценка ум. способ.	50	59,76000	60,00000	0,108701	0,336601	0,446733	0,661908

В приведенных результатах видно, что: средняя арифметическая и медиана значительно не отличаются, при этом показатель асимметрии не превышает трех ошибок асимметрии, а показатель эксцесса не превышает трех ошибок эксцесса. Таким образом, соблюдение этих условий позволяет сделать вывод о нормальном характере распределения данных методики изучения самооценки.

3.2. Выбор статистического критерия различия

Одной из наиболее часто встречающихся статистических задач, с которыми сталкивается психолог, является задача сравнения результатов обследования какого-либо психологического признака в разных условиях измерения (например, «до» и «после» определенного воздействия) или обследования контрольной и экспериментальной групп. Помимо этого нередко возникает необходимость оценить характер изменения того или иного психологического показателя в одной или нескольких группах в разные периоды времени или выявить динамику изменения этого показателя под влиянием экспериментальных воздействий.

Для решения подобных задач используется достаточно большой набор статистических способов, называемых в наиболее общем виде критериями различий. Эти критерии позволяют оценить степень статистической достоверности различий между разнообразными показателями, измеренными согласно плану проведения психологического исследования.

Существует достаточно большое количество критериев различий. Каждый из них имеет свою специфику, различаясь между собой по различным основаниям. Одним из таких оснований является тип измерительной шкалы, для которой предназначен тот или иной критерий. Например, с помощью некоторых критериев можно обрабатывать данные, полученные только в номинальных шкалах. Ряд критериев дает возможность обрабатывать данные, полученные в порядковой, интервальной и шкале отношений.

Критерии различаются также по максимальному объему выборки, который они могут охватить, а также и по количеству выборок, которые можно сравнивать между собой с их помощью. Так, существуют критерии, позволяющие оценить различия сразу в трех и большем числе выборок. Некоторые критерии позволяют сопоставлять неравные по численности выборки.

Еще одним признаком, дифференцирующим критерии, служит само качество выборки: она может быть связанной (зависимой) или несвязанной (независимой). Выборки также могут быть взяты из одной или нескольких генеральных совокупностей. Именно эта характеристика выборки служит наиболее важным основанием, по которому выбираются критерии.

Кроме того, критерии различаются по мощности. Психолог может решать экспериментальные задачи с использованием разных статистических критериев. При этом возможна такая ситуация, что один критерий позволяет обнаружить различия, а другой критерий различий не выявляет. Последнее означает, что первый критерий оказывается более мощным, чем другой.

Все критерии различий условно подразделены на две группы: параметрические и непараметрические критерии.

Критерий различия называют параметрическим, если он основан на конкретном типе распределения генеральной совокупности (как правило, нормальном) или использует параметры этой совокупности (средние, дисперсии и т. д.). Критерий различия называют непараметрическим, если он не базируется на предположении о типе распределения генеральной совокупности и не использует параметры этой совокупности [2].

При нормальном распределении генеральной совокупности параметрические критерии обладают большей мощностью по сравнению с непараметрическими. Иными словами, они способны с большей достоверностью отвергать нулевую гипотезу, если последняя неверна. По этой причине в тех случаях, когда выборки взяты из нормально распределенных генеральных совокупностей, следует отдавать предпочтение параметрическим критериям. Однако, как показывает практика, подавляющее большинство данных, получаемых в психологических экспериментах, не распределены нормально [2]. Поэтому применение параметрических критериев при анализе результатов психологических исследований может привести к ошибкам в статистических выводах. В таких случаях непара-

метрические критерии оказываются более мощными, т. е. способными с большей достоверностью отвергать нулевую гипотезу.

Анализ частотных распределений эмпирических данных, полученных в ходе психодиагностических обследований испытуемых, необходим для выбора способов их математико-статистической обработки. От результатов такой оценки во многом зависит, например, выбор критериев выявления различий в психологических характеристиках, полученных в различных группах испытуемых, или сдвига исследуемой психологической характеристики в ходе экспериментального воздействия или формирующего эксперимента.

Для правильной статистической обработки результатов психологических исследований должны соблюдаться правила по выбору параметрических и непараметрических критериев, позволяющих выявлять различия в исследуемых психологических характеристиках.

Для психологических характеристик, имеющих нормальное распределение или близкое к нормальному, необходимо использовать параметрические критерии, которые являются более мощными, чем непараметрические критерии. Достоинством непараметрических критериев является то, что они позволяют проверять статистические гипотезы независимо от формы распределения.

В подразделе 2.1 представлен способ оценки нормальности распределения выборочных данных на основе первичных статистик.

В данном разделе приводится пример более точной и строгой оценки нормальности распределения с помощью критерия «хи-квадрат» Пирсона $\chi^2_{\text{эмп}}$.

Критерий «хи-квадрат» в одном из вариантов своего использования применяется при расчете согласия эмпирического и предполагаемого теоретического распределения – в этом случае проверяется гипотеза H_0 об отсутствии различий между теоретическим и эмпирическим распределениями [2].

Рассмотрим задачу, в которой в качестве теоретического будет использоваться нормальное распределение (задача взята из учебника О. Ю. Ермолаева «Математическая статистика для психологов», 2002) [2].

Задача. У 267 человек был измерен рост. Вопрос состоит в том, будет ли полученное в этой выборке распределение роста близко к нормальному?

Решение. Измерения проводились с точностью до 0,1 см, и все полученные величины роста оказались в диапазоне от 156,5 до 183,5. Для расчета по критерию «хи-квадрат» целесообразно разбить этот диапазон на интервалы, например по 3 см каждый. Тогда все экспериментальные данные будут распределены по 9 интервалам $(183,5-156,5)/3=9$. При этом центрами интервалов будут следующие числа: 158, 161, 164 и т. д. до 182.

При измерении роста в каждый из этих интервалов попало какое-то количество людей – эта величина для каждого интервала и будет эмпирической частотой, обозначаемой в дальнейшем как $f_{эi}$.

Чтобы применить расчетную формулу для вычисления значения критерия «хи-квадрат»:

$$\chi^2_{y\ddot{u}} = \sum_{i=1}^k \Delta_i^2 / f_{in} , \quad (3.1)$$

где Δ_i – разность между эмпирическими и «теоретическими» частотами;

k – количество разрядов признака;

f_{mi} – вычисленная, или «теоретическая» частота.

Прежде всего, необходимо вычислить теоретические частоты. Для этого по всем выборочным данным нужно вычислить среднее M и стандартное отклонение σ по известным из литературных источников формулам [7]. В качестве тренировки слушателям предлагается проверить себя и самостоятельно рассчитать обозначенные параметры распределения выборочных данных двумя способами: вручную и с помощью программ статистической обработки данных; сравнить полученные двумя способами значения одноименных параметров между собой, а также со значениями, указанными в данном учебном пособии.

Для наших выборочных данных величина M оказалась равной 166,22 и величина σ получилась равной 4,06.

Затем для каждого выделенного интервала следует подсчитать величины o_i (индекс i меняется от 1 до 9) по формуле:

$$o_i = (x_i - M) / \sigma. \quad (3.2)$$

Величины o_i называются **нормированными частотами**. Удобнее производить их расчет в приведенной ниже табл. 3.1.

Таблица 3.1

№ 1	№ 2	№ 3	№ 4	№ 5
Центры интервалов x_i	Эмпирические частоты f_{xi}	Нормированные частоты o_i	Ординаты нормальной кривой $f(o_i)$	Расчетные теоретические частоты f_{mi}
158	3	-2,77	0,0086	1,6
161	9	-2,03	0,0508	10,0
164	31	-1,29	0,1736	34,3
167	71	-0,55	0,3429	67,8
170	82	+0,19	0,3918	77,6
173	46	+0,93	0,2589	51,2
176	19	+1,67	0,0989	19,5
179	5	+2,41	0,0219	4,4
182	1	+3,15	0,0028	0,6
S	267	-	-	267

Затем для каждой o_i по табл. 2 (см. приложение 1) находят величины $f(o_i)$, которые называются ординаты нормальной кривой.

Расчет теоретических частот осуществляется для каждого интервала по следующей формуле:

$$f_{ti} = f(o_i) \times (n \times \lambda) / \sigma, \quad (3.3)$$

где $n=267$ – общая величина выборки;

$\lambda=3$ – величина интервала;

σ – стандартное отклонение.

Все рассчитанные значения заносятся в соответствующие столбцы табл. 3.1. После ее заполнения все готово для работы с критерием «хи-квадрат» на основе стандартной таблицы. В целях упрощения расчетов сократим число интервалов до семи за счет сложения двух верхних и двух нижних частот. Тогда стандартная таблица для вычисления «хи-квадрат» будет выглядеть следующим образом.

В случае оценки равенства эмпирического распределения нормальному число степеней свободы определяется особым образом: из общего числа интервалов вычитается число 3. Таким образом, число степеней свободы в нашем примере будет равно $\nu=4$.

Таблица 3.2

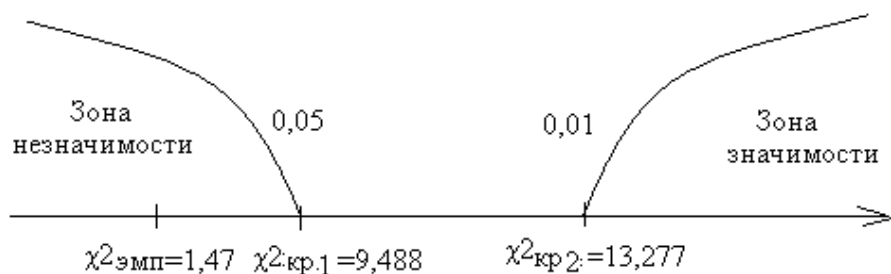
№ 1	№ 2	№ 3	№ 4	№ 5	№ 6
Номера интервалов (альтернативы)	$f_{\text{эi}}$	$f_{\text{тi}}$	$(f_{\text{эi}} - f_{\text{тi}})$	$(f_{\text{эi}} - f_{\text{тi}})^2$	$(f_{\text{эi}} - f_{\text{тi}})^2 / f_{\text{тi}}$
1	12	11,6	+0,4	0,16	0,01
2	31	34,3	-3,3	10,89	0,32
3	71	67,8	+3,2	10,24	0,15
4	82	77,6	+4,4	19,36	0,25
5	46	51,2	-5,2	27,04	0,53
6	19	19,5	-0,5	0,25	0,01
7	6	5,0	+1,0	1,00	0,20
S	267	267	0		$\chi^2_{\text{эмп}} = 1,47$

По табл. 3 (см. приложение 1) находим критические значения критерия «хи-квадрат»:

$$\chi^2_{\text{кр.1}} = 9,488 \text{ (для } p \leq 0,05),$$

$$\chi^2_{\text{кр.2}} = 13,277 \text{ (для } p \leq 0,01).$$

Строим «ось значимости»:



Полученная величина эмпирического значения «хи-квадрат» попала в зону незначимости, поэтому необходимо принять гипотезу H_0 об отсутствии различий. Следовательно, существуют все основания утверждать, что наше эмпирическое распределение близко к нормальному.

Несмотря на некоторую «громоздкость» вычислительных процедур, этот способ расчета дает наиболее точную оценку совпадения эмпирического и нормального распределений.

Пример применения в психологических исследованиях критерия «хи-квадрат» приведен в приложении 3 настоящего учебного пособия (пример 2).

Домашнее задание. В качестве домашнего задания для тех слушателей, которые хорошо владеют навыками статистической обработки данных с помощью компьютерных программ, предлагается рассчитать критерий «хи-квадрат» на одном из известных и популярных статистических пакетов и сравнить с результатом решения задачи, полученным вручную, как описано в настоящем учебном пособии.

Возвращаясь к разобранный задаче, не лишним будет обозначить критерии выявления различий, которыми целесообразно пользоваться при обработке эмпирических данных. При принятии гипотезы H_0 для обработки

данных выбирают t -критерий Стьюдента или F -критерий Фишера. При опровержении гипотезы H_0 выбирают критерий знаков G , парный критерий T -Вилкоксона, критерий Фридмана, критерий тенденций Пейджа, критерий Макнамары, которые рассматриваются для условий связанных выборок, или критерий U Вилкоксона-Манна-Уитни, критерий Q Розенбаума, H -критерий Крускала-Уоллиса, S -критерий тенденций Джонкира, которые рассматриваются для условий несвязанных выборок. Все эти критерии подробно, с примерами рассматриваются у О. Ю. Ермолаева, Е. В. Сидоренко [2, 8].

Одной из наиболее часто встречающихся задач при обработке данных является оценка достоверности отличий между двумя или более рядами значений [4]. В математической статистике существует ряд способов для этого. Для использования большинства мощных критериев требуются дополнительные вычисления, обычно весьма развернутые.

Перед психологом часто встает задача оценки достоверности различий, используя ранее вычисленные статистики. При сравнении средних значений признака говорят о достоверности (недостоверности) отличий средних арифметических, а при сравнении изменчивости показателей – о достоверности (недостоверности) отклонений сигм (дисперсий) и коэффициентов вариации.

Достоверность различий средних арифметических можно оценить по достаточно эффективному параметрическому **критерию Стьюдента**. Он вычисляется по формуле:

$$t = \frac{M_1 - M_2}{(m_1^2 + m_2^2)^{1/2}}, \quad (3.4)$$

где M_1 и M_2 – значения сравниваемых средних арифметических;

m_1 и m_2 – соответствующие величины статистических ошибок средних арифметических (вычисляются по формуле 3.4).

Значения критерия Стьюдента t для трех уровней значимости (p) приведены в табл. 4 (см. приложение 1). Число степеней свободы определяется по формуле:

$$d = n_1 + n_2 - 2, \quad (3.5)$$

где n_1 и n_2 – объемы сравниваемых выборок.

С уменьшением объемов выборок ($n < 10$) критерий Стьюдента становится чувствительным к форме распределения исследуемого признака в генеральной совокупности. Поэтому в сомнительных случаях рекомендуется использовать непараметрические методы или сравнивать полученные значения с критическими (приведенными в таблице) для более высокого уровня значимости.

Решение о достоверности различий принимается в том случае, если вычисленная величина t превышает табличное значение для данного числа степеней свободы. В тексте публикации или научного отчета указывают наиболее высокий уровень значимости из трех: 0,05, 0,01 или 0,001.

Если превышены 0.05 и 0.01, то пишут (обычно в скобках) $p = 0,01$ или $p < 0,01$. Это означает, что оцениваемые различия все же случайны только с вероятностью не более 1 из 100 шансов. Если превышены табличные значения для всех трех уровней: 0,05, 0,01 и 0,001, то указывают $p = 0,001$ или $p < 0,001$, что означает случайность выявленных различий между средними не более одного из 1000 шансов.

Пример оценки достоверности различий средних арифметических по критерию Стьюдента:

$$M_1 = 113,3, m_1 = 2,4, n_1 = 13 ; M_2 = 103,3, m_2 = 2,6, n_2 = 16.$$

$$t = \frac{113,3 - 103,3}{(2,4^2 + 2,6^2)^{1/2}} = 2,8 .$$

Для числа степеней свободы $d = 13 + 16 - 2 = 27$ вычисленная величина превышает табличную 2,77 для вероятности $p=0,01$. Следовательно, различия между средними достоверны на уровне 0,01.

Приведенная формула проста. Используя ее, можно с помощью простейшего бытового калькулятора с памятью вычислить t критерий без промежуточных записей.

Следует помнить, что при любом численном значении критерия достоверности различия между средними этот показатель оценивает не степень выявленного различия (она оценивается по самой разности между средними), а лишь статистическую достоверность его, т. е. право распространять полученный на основе сопоставления выборок вывод о наличии разницы на все явление (весь процесс) в целом. Низкий вычисленный критерий различия не может служить доказательством отсутствия различия между двумя признаками (явлениями), ибо его значимость (степень вероятности) зависит не только от величины средних, но и от численности сравниваемых выборок. Он говорит не об отсутствии различия, а о том, что при данной величине выборок оно статистически недостоверно: слишком велик шанс, что разница при данных условиях определения случайна, слишком мала вероятность ее достоверности.

Степень, т. е. величину выявленного различия, желательно оценивать, опираясь на содержательные критерии. Вместе с тем, для психологического исследования весьма характерно наличие множества показателей, которые, по существу, являются условными баллами, и валидность оценивания с помощью них еще предстоит доказать. Чтобы избежать большей произвольности, в таких случаях также приходится опираться на статистические параметры.

Пожалуй, наиболее распространено для этого использование сигмы. Разницу между двумя значениями в одну сигму и более можно считать достаточно выраженной. Если сигма подсчитана для ряда значений более 35, то достаточно выраженной можно рассматривать и разницу в 0,5

сигмы. Однако для ответственных выводов о том, насколько велика разница между значениями, лучше использовать строгие критерии.

При оценке различий в распределениях, далеких от нормального, непараметрические критерии могут выявить значимые различия, в то время как параметрические критерии таких различий не обнаружат. Важно отметить, что непараметрические критерии выявляют значимые различия и в том случае, если распределение близко к нормальному.

Рекомендации к выбору критерия различий.

При подготовке экспериментального исследования психолог должен заранее запланировать характеристики сопоставляемых выборок (прежде всего, связность – несвязность и однородность), их величину (объем), тип измерительной шкалы и вид используемого критерия различий. Последовательно это можно представить в виде следующих этапов [2]:

1. Прежде всего, следует определить, является ли выборка связной (зависимой) или несвязной (независимой).

2. Следует определить однородность–неоднородность выборки.

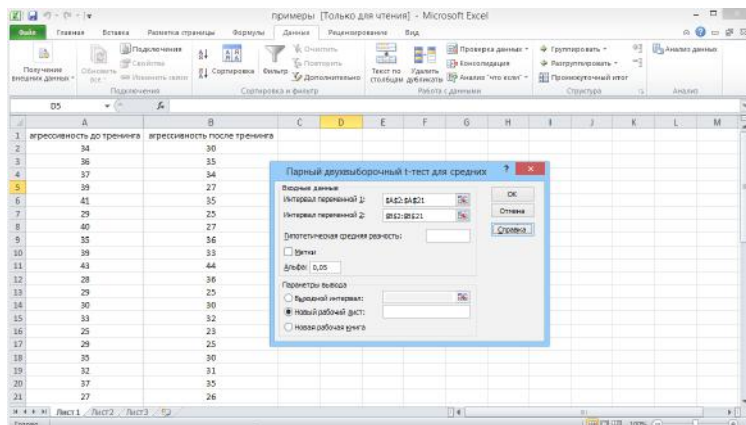
3. Затем следует оценить объем выборки и, зная ограничения каждого критерия по объему, выбрать соответствующий критерий.

4. Если используемый критерий не выявил различия, следует применить более мощный критерий.

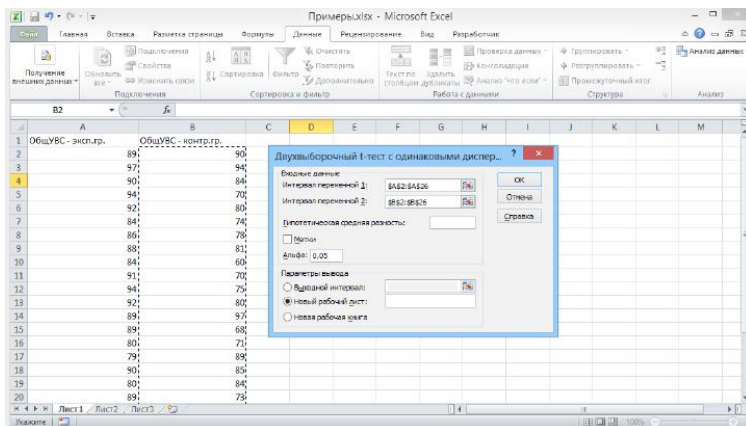
5. Если в распоряжении психолога имеется несколько критериев, то следует выбирать те из них, которые наиболее полно используют информацию, содержащуюся в экспериментальных данных.

6. При малом объеме выборки следует увеличивать величину уровня значимости (не менее 1 %, т. е. $p=0,01$), так как небольшая выборка и низкий уровень значимости приводят к увеличению вероятности принятия ошибочных решений.

Для решения поставленной задачи методами статистического анализа Excel на панели инструментов нужно выбрать «Данные» > «Анализ данных» > «Описательная статистика» > «Парный двухвыборочный t-тест для средних»:



Столбцы документа содержат «сырые» данные тестирования, полученные «до» и «после» проведения тренинга. Перед обработкой данных необходимо определить, соблюдены ли все условия применения t -критерия Стьюдента и обозначить в строках интервала первой и второй переменной массив данных агрессивности «до» и «после» проведения тренинга:



В нашем примере статистический пакет Excel произвел следующие расчеты:

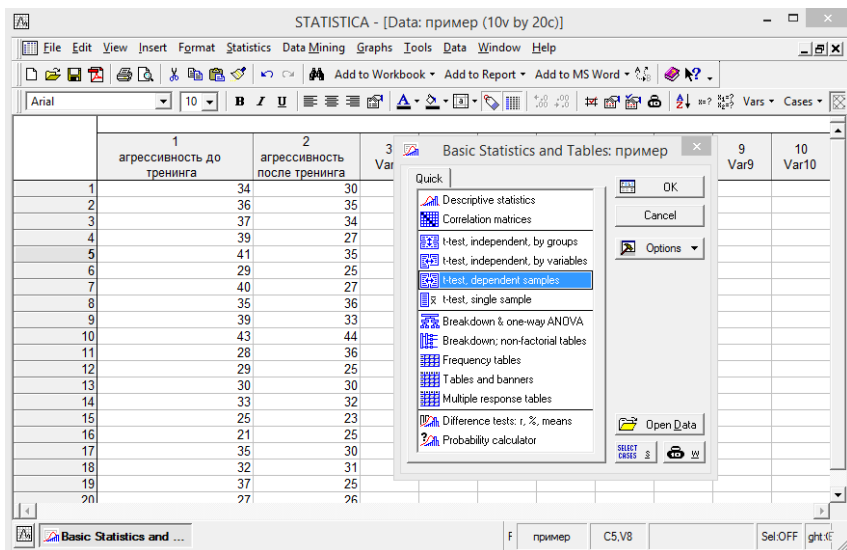
Парный двухвыборочный t-тест для средних		
	Переменная 1	Переменная 2
Среднее	33,9	30,95
Дисперсия	26,41052632	26,99736842
Наблюдения	20	20
Корреляция Пирсона	0,618711901	
Гипотетическая разность средних	0	
df	19	
t-статистика	2,923387521	
P(T<=t) одностороннее	0,004359245	
t критическое одностороннее	1,729132812	
P(T<=t) двухстороннее	0,00871849	
t критическое двухстороннее	2,093024054	

Таким образом, можно увидеть, что эмпирическое значение t -критерия равно 2,92, а показатель p -level («Р двухстороннее»), который является самым важным в данной таблице, гораздо меньше 0,05 (в нашем примере 0,00871849), при этом среднее значение первой группы больше, чем у второй, что позволяет нам принять гипотезу о статистически достоверном снижении агрессии, произошедшем в группе после тренинга.

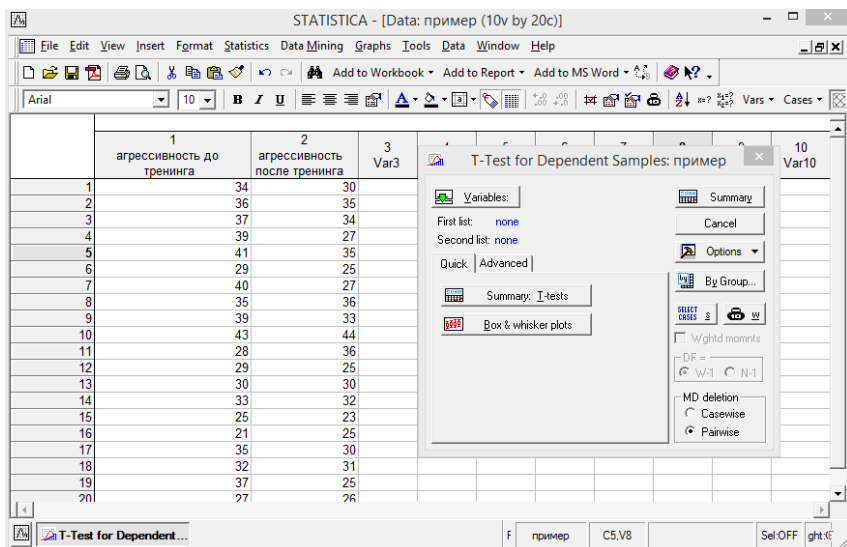
Решение в Statistica. После проведения подготовительных работ с документом (см. приложение 4) он будет выглядеть следующим образом:

	1 агрессивность до тренинга	2 агрессивность после тренинга	3 Var3	4 Var4	5 Var5	6 Var6	7 Var7	8 Var8	9 Var9	10 Var10
1	34	30								
2	36	35								
3	37	34								
4	39	27								
5	41	35								
6	29	25								
7	40	27								
8	35	36								
9	39	33								
10	43	44								
11	28	36								
12	29	25								
13	30	30								
14	33	32								
15	25	23								
16	21	25								
17	35	30								
18	32	31								
19	37	25								
20	27	26								

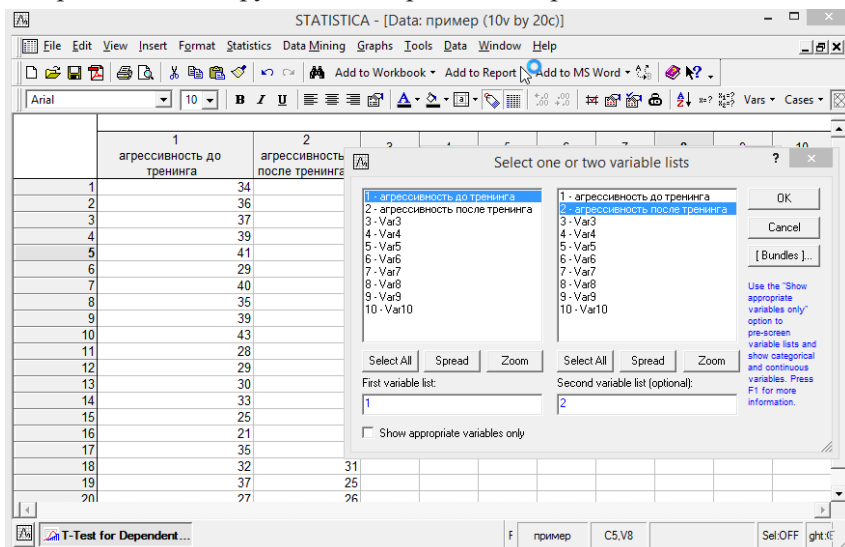
Выберите «Statistics» («Статистика») > «Basic statistics/Tables» («Основные статистики и таблицы») > «t-test, dependent samples» («t-тест для связанных (зависимых) выборок»):



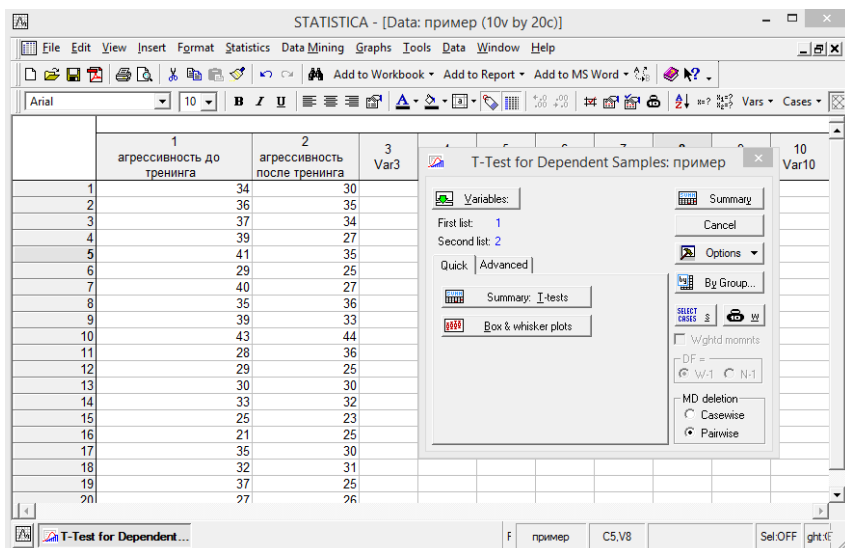
В окне-меню дважды нажмите на строчку «t-test, dependent samples»:



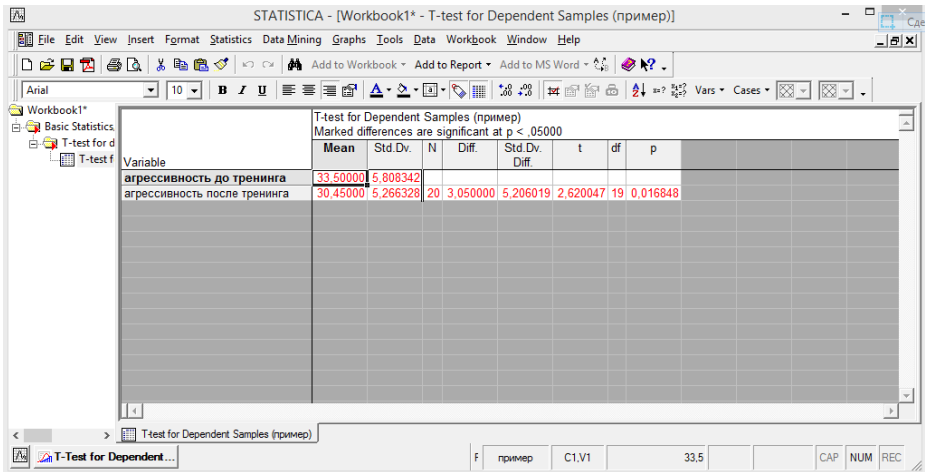
Выберите две сравниваемые группы переменных. В нашем примере в первую группу отнесите значения агрессивности у испытуемых до проведения тренинга, во вторую – после проведения тренинга. Нажмите на «ОК»:



В появившемся окне нажмите на функцию «Summary» («Итог»):



В итоговой статистической таблице выделенные красным цветом показатели свидетельствуют о достоверности различий между значениями двух сравниваемых групп. Показатель t-критерия составил в нашем примере 2,62 при уровне значимости (p-level) 0,017, а среднее значение первой группы больше, чем у второй (33,5 и 30,45 соответственно). Таким образом, результаты проведенного анализа показали, что после проведенного тренинга у двадцати обучающихся статистически достоверно выявлено снижение агрессивности:

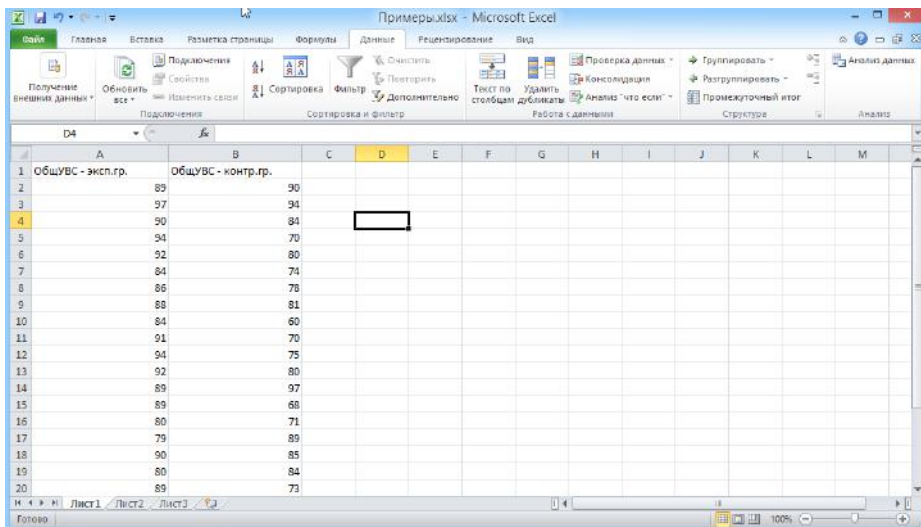


Далее рассмотрим применение *t*-критерия на примере несвязанных выборок, когда в целях эксперимента для сравнения привлекаются данные двух или более выборок, причем эти выборки могут быть взяты из одной или из разных генеральных совокупностей. Таким образом, для несвязанных выборок характерно, что в них обязательно входят разные испытуемые.

Пример задачи. Исследователь должен выяснить, отличаются ли показатели общего уровня волевой самооценки (ОбщУВС) в двух группах испытуемых: экспериментальной, члены которой участвовали в

тренинге развития волевых качеств, и контрольной – не участвовали в подобном тренинге. Численность обучающихся в каждой группе 25 человек.

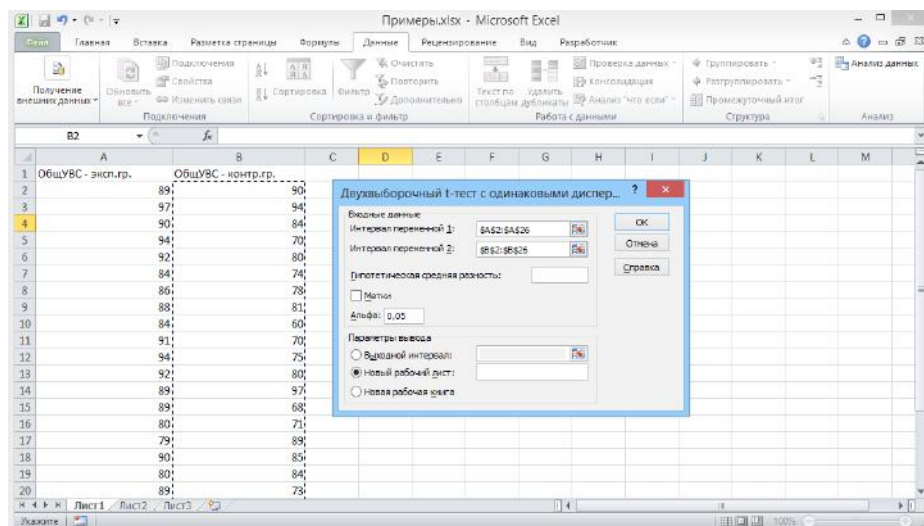
Решение в Excel. После проведения подготовительных работ с документом он будет выглядеть следующим образом:



	A	B	C	D	E	F	G	H	I	J	K	L	M
1	ОбщУВС - эксп.гр.	ОбщУВС - контр.гр.											
2		89	90										
3		97	94										
4		90	84										
5		94	70										
6		92	80										
7		84	74										
8		86	78										
9		88	81										
10		84	60										
11		91	70										
12		94	75										
13		92	80										
14		89	97										
15		89	68										
16		80	71										
17		79	89										
18		90	85										
19		80	84										
20		89	73										

В первом столбце размещены данные экспериментальной группы, во втором столбце – контрольной.

Для решения поставленной задачи методами статистического анализа Excel на панели инструментов нужно выбрать «Данные» > «Анализ данных» > «Описательная статистика» > «Парный выборочный t -тест для средних»: двухвыборочный t -тест с одинаковыми дисперсиями. Перед обработкой данных необходимо определить, соблюдены ли все условия применения t -критерия Стьюдента с одинаковыми дисперсиями – эмпирические данные должны быть из одной и той же генеральной совокупности. Далее необходимо обозначить в строках интервала массив данных сравниваемого показателя (ОбщУВС) в экспериментальной и контрольной группах:



В нашем примере статистический пакет Excel произвел следующие расчеты:

Двухвыборочный t-тест с одинаковыми дисперсиями		
	Переменная 1	Переменная 2
Среднее	88	78,4
Дисперсия	22,25	87,16666667
Наблюдения	25	25
Объединенная дисперсия	54,70833333	
Гипотетическая разность средних	0	
df	48	
t-статистика	4,588803885	
P(T<=t) одностороннее	1,60962E-05	
t критическое одностороннее	1,677224196	
P(T<=t) двухстороннее	3,21924E-05	
t критическое двухстороннее	2,010634758	

Таким образом, можно увидеть, что эмпирическое значение t -критерия равно – 4,59, а показатель p -level («P двухстороннее»), гораз-

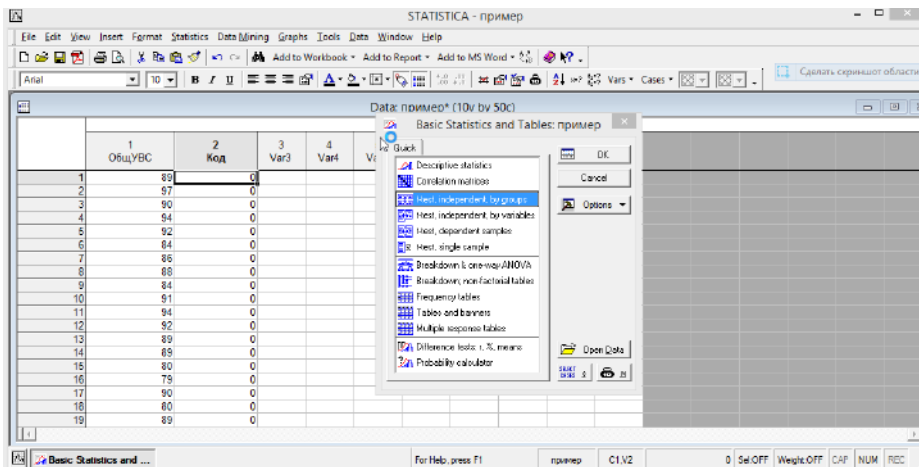
до меньше 0,05 (в нашем примере 0,0000321924), при этом среднее значение первой группы больше, чем у второй, что позволяет нам принять гипотезу о статистически достоверных различиях общего уровня волевой самооценки испытуемых в двух группах сравнения.

Решение в Statistica. После проведения подготовительных работ с документом (см. приложение 4) он будет выглядеть следующим образом:

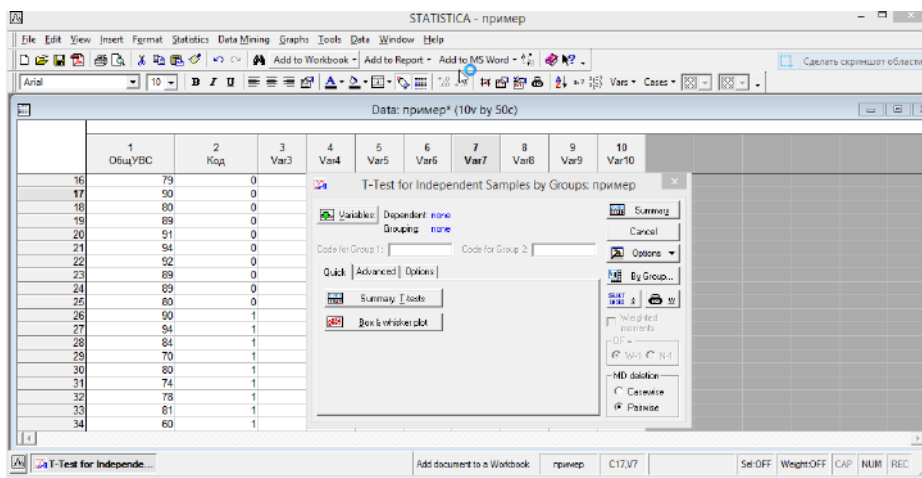
	1 ОбщУВС	2 Код	3 Var3	4 Var4	5 Var5	6 Var6	7 Var7	8 Var8	9 Var9	10 Var10
16	79	0								
17	90	0								
18	80	0								
19	89	0								
20	91	0								
21	94	0								
22	92	0								
23	89	0								
24	89	0								
25	80	0								
26	90	1								
27	94	1								
28	84	1								
29	70	1								
30	80	1								
31	74	1								
32	76	1								
33	81	1								
34	60	1								

В первом столбце («ОбщУВС») размещены показатели общего уровня волевой самооценки всех испытуемых, во втором столбце («Код») напротив каждого показателя проставлен код принадлежности испытуемого к одной из двух исследовательских групп. Значения кодов могут быть любыми, главное, чтобы они не совпадали. Для простоты изложения в нашем примере код экспериментальной группы равен нулю («0»), контрольной – единице («1»). Таким образом, показатели испытуемых первой группы соотносятся с кодом «0», а второй – с кодом «1».

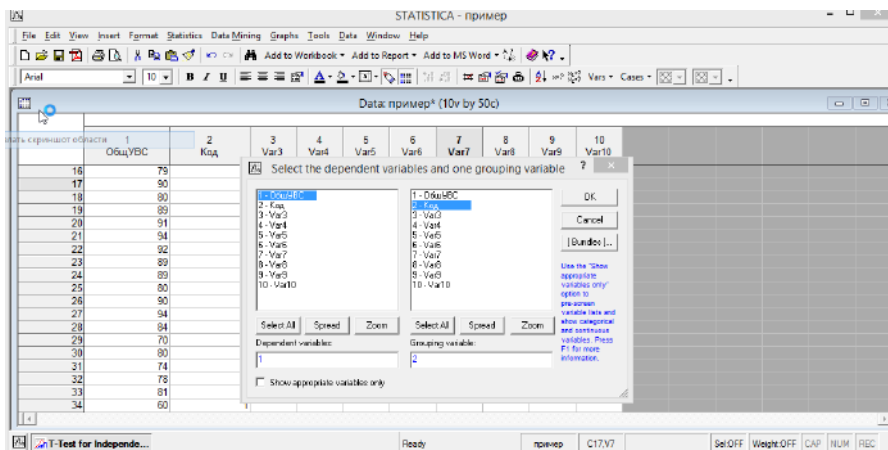
Выберите «Statistics» («Статистика») > «Basic statistics/Tables» («Основные статистики и таблицы») > «t-test, independent, by groups» («t-тест для несвязанных (независимых) выборок»):



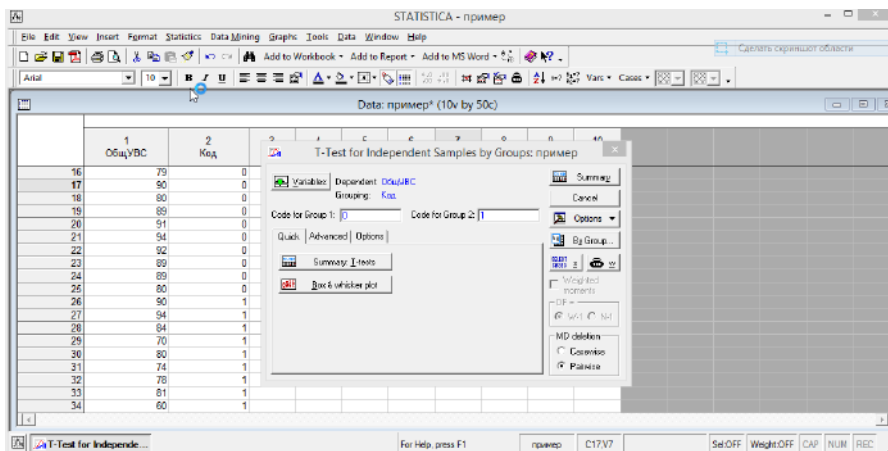
В окне-меню дважды нажмите на строчку «t-test, independent, by groups»:



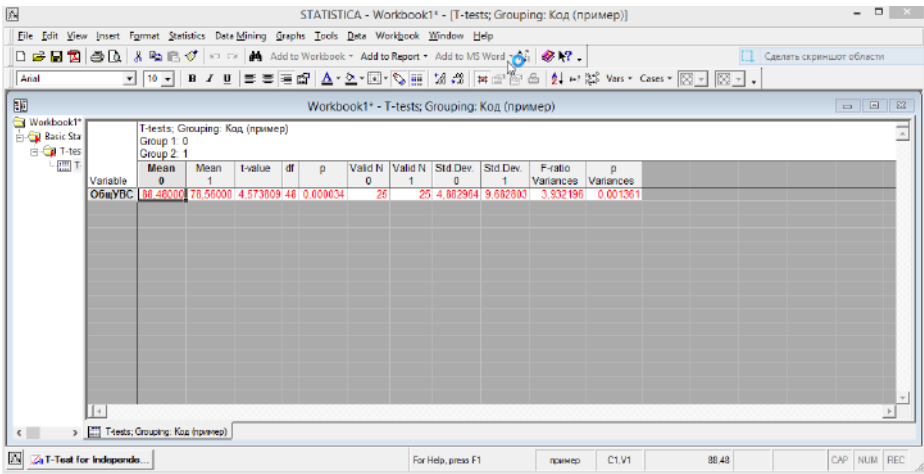
В появившемся окне функций «t-test, independent, by groups» нажмите на «Variables» («Переменные»):



Выберите две группы переменных. В левой нижней части таблицы необходимо обозначить зависимую переменную (dependent variable), которая в нашем примере связана с общим уровнем волевой самооценки, поэтому нажмите на «1-ОбщУБС». В правой нижней части таблицы необходимо обозначить независимую переменную (grouping variable), которая в нашем примере связана с экспериментальной и контрольной группами, выделенных кодами «0» и «1», поэтому нажмите на «2-Код». После окончания работы с данным окном нажмите на «ОК»:



В новом появившемся окне вы увидите, что программа все значения испытуемых первой группы связала с кодом «0», а значения испытуемых второй группы - с кодом «1». Нажмите на «Summary»:



В итоговой статистической таблице выделенные красным цветом показатели свидетельствуют о достоверности различий между значениями двух сравниваемых групп. Показатель t-критерия составил в нашем примере 4,57 при уровне значимости (p-level) 0,00034, при этом среднее значение первой группы больше, чем у второй.

Таким образом, результаты проведенного анализа показали, что показатели общего уровня волевой самооценки (ОбщУВС) у испытуемых экспериментальной группы, члены которой участвовали в тренинге развития волевых качеств, статистически достоверно выше, чем у сверстников из контрольной группы.

Контрольные вопросы и задания

1. Проведите проверку на нормальность распределения данных, полученных с помощью любой психодиагностической методики. Ответьте на вопрос: «Зависит ли результат подобной проверки от размера выборочных данных?»

2. По какому принципу выбирается статистический критерий различия выборочных данных?

3. Какие особенности выявления различий исследуемых данных с помощью критериев: t -критерия Стьюдента и «хи-квадрат» Пирсона $\chi^2_{\text{эмп}}$?

4. Каких рекомендаций следует придерживаться исследователю при выборе критериев различий?

5. Подготовьте данные в электронных таблицах по результатам проведенных исследований с помощью психодиагностических методик для выявления различий по t -критерию Стьюдента для связанных и несвязанных выборок.

6. Проведите обработку данных, направленную на выявление различий с помощью t -критерия Стьюдента, используя табличный редактор Excel в лицензионном пакете Microsoft Office и программу Statistica, с учетом связности или несвязности выборочных данных.

4. Корреляционный анализ

В практике экспериментальных исследований нередко случаи, когда предполагается наличие связанных изменений каких-либо двух статистических признаков. Например, представляются взаимозависимыми вариации величины роста и веса тела людей (прямая связь), силы мышц и их подвижности (обратная связь) и т. д. Такого рода связи и закономерности не являются строго однозначными или функциональными, они, так же как и сами вариации признаков, являются статистическими, или корреляционными.

Теория корреляционного исследования, основанная на представлениях о мерах корреляционной связи, разработана К. Пирсоном.

Корреляционным называется исследование, проводимое для подтверждения или опровержения гипотезы о статистической связи между несколькими (двумя и более) переменными. В психологии в качестве переменных могут выступать психические свойства, процессы, состояния и др. [2].

Корреляция – это связь между статистическими вариациями (выборками) по различным признакам, между влияниями каких-либо двух факторов, формирующих данное статистическое распределение.

Корреляция в прямом переводе означает «соотношение». Если изменение одной переменной сопровождается изменением другой, то можно говорить о корреляции этих переменных. Наличие корреляции двух переменных ничего не говорит о причинно-следственных зависимостях между ними, но дает возможность выдвинуть такую гипотезу. Отсутствие же корреляции позволяет отвергнуть гипотезу о причинно-следственной связи переменных [3].

Коэффициент корреляции – это математический показатель силы (тесноты) связи между двумя сопоставляемыми статистическими признаками.

По какой бы формуле ни вычислялся коэффициент корреляции, его величина колеблется в пределах от -1 до +1. Смысл крайних значений коэффициента состоит в следующем:

- если коэффициент корреляции равен $+1$, значит, связь между признаками однозначна (функциональная, нестатистическая) по типу прямо пропорциональной зависимости;

- если коэффициент равен -1 , то связь также является функциональной, но по типу обратной пропорциональности;

- нулевая величина коэффициента корреляции говорит о полном отсутствии связи (по типу линейной) между сопоставляемыми признаками.

Психологам часто хочется получать ответы на такие вопросы, эмпирические данные для которых могут быть собраны только в корреляционном исследовании. Проверка соответствующих гипотез, если они понимаются именно как гипотезы о взаимосвязях переменных, а не о причинной зависимости, может вести к обоснованным выводам.

Всякое вычисленное (эмпирическое) значение коэффициента корреляции должно быть проверено на статистическую значимость (см. табл. 1 или 2 приложения 3).

Если эмпирическое значение коэффициента корреляции меньше или равно табличному для $p=0,05$, то корреляция является незначимой. Если вычисленное значение коэффициента корреляции больше табличного для $p=0,01$, корреляция статистически значима (существенна, реальна). В случае, когда величина коэффициента заключена между двумя табличными, на практике говорят о значимости корреляции для $p=0,05$. Однако строго вероятностная трактовка этого факта несколько иная: мы не можем утверждать отсутствия корреляции, но ее статистически доказанного наличия также еще нет.

Простейшей формой коэффициента корреляции является *коэффициент ранговой корреляции* R_s (коэффициент Спирмена), который измеряет связь между рангами (местами) данной варианты по разным признакам, но не между собственными величинами варианты. Здесь исследуется связь качественная, чем строго количественная, хотя ранг сам по себе – это уже и количественный признак.

$$R_s = 1 - 6 \times \Sigma d^2 / (n^3 - n), \quad (4.1)$$

где n – объем совокупности, длина одного статистического ряда;

d – разность между рангами каждой варианты по двум коррелируемым признакам.

Пример вычисления. Десять испытуемых (А, Б, В и т. д.) расположились в порядке увеличения возраста и пространственного порога в следующих последовательностях (табл. 4.1):

Таблица 4.1

Испытуемые	Ранг по возрасту	Ранг по пространств. порогу	d	d ²
А	1	6	-5	25
Б	2	5	-3	9
В	3	2	1	1
Г	4	1	3	9
Д	5	10	-5	25
Е	6	4	2	4
Ж	7	9	-2	4
З	8	7	1	1
И	9	8	1	1
К	10	3	7	49

$n = 10$

$\Sigma d^2 = 128$

$$R_s = 1 - 6 \times 128 / (1000 - 10) = 1 - 768 / 990 = 0,22,$$

$$R_s = 0,22.$$

Так как по данным табл. 4.1 (см. приложение 2) $R_{s,0,05} = 0,64$, и эмпирическое значение $R_s < R_{s,0,05}$, корреляция между местами испытуемых по величине порогов и по возрасту не является статистически значимой.

Пример применения в психологических исследованиях коэффициента корреляции R_s Спирмена приведен в приложении 3 настоящего учебного пособия (пример 1).

Используемый в исследованиях коэффициент корреляции рангов, предложенный К. Спирменом, относится к непараметрическим показателям связи между переменными, измеренными в ранговой шкале. При расчете этого коэффициента не требуется никаких предположений о характере распределений признаков в генеральной совокупности. Данный коэффициент определяет степень тесноты связи порядковых признаков, которые в этом случае представляют собой ранги сравниваемых величин [6].

Величина коэффициента линейной корреляции Спирмена лежит в интервале +1 и -1. Он может быть положительным и отрицательным, характеризуя направленность связи между двумя признаками, измеренными в ранговой шкале.

Для применения коэффициента корреляции Спирмена (r) необходимо соблюдать следующие условия [2]:

1. Сравнимые переменные должны быть получены в порядковой (ранговой) шкале, но могут быть измерены также в шкале интервалов и шкале отношений.

2. Число варьирующих признаков в сравниваемых переменных X и Y должно быть одинаковым.

Пример задачи. Рассмотрим, как при помощи статистического пакета «MS Office Excel» провести расчет коэффициента корреляции Пирсона, используемого на массиве эмпирических данных, подпадающих под нормальное распределение.

Исследователь должен выяснить, как связаны между собой значения пяти шкал методики изучения самооценки (Дембо-Рубинштейн): 1-«ум», 2-«характер», 3-«красота», 4-«здоровье», 5-«счастье» у 25 испытуемых.

Решение в Excel. После проведения подготовительных работ с документом он будет выглядеть следующим образом:

Microsoft Excel - примеры [Только для чтения] - Microsoft Excel

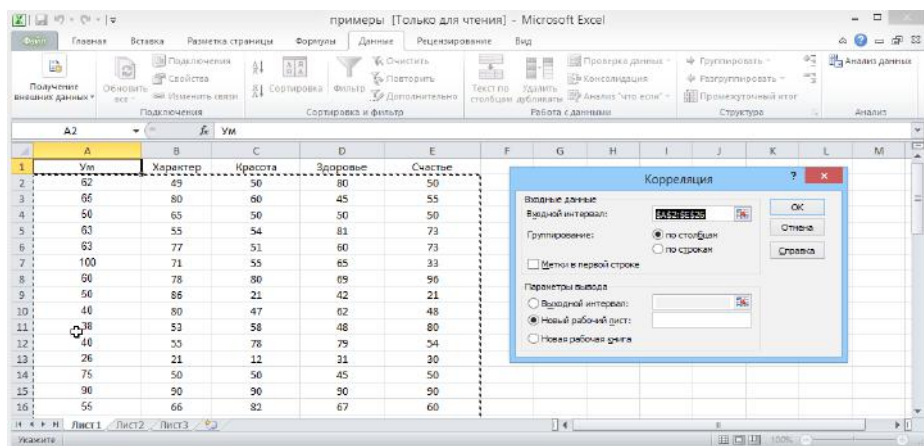
Файл | Главная | Вставка | Разметка страницы | Формулы | Данные | Рецензирование | Вид

Получение внешних данных | Обновить все | Подключения | Свойства | Изменить связи | Подключения | Сортировка | Фильтр | Дополнительно | Проверка данных | Консолидация | Анализ "что если" | Группировать | Разгруппировать | Промежуточный итог | Анализ данных

	A	B	C	D	E	F	G	H	I	J	K	L	M
	Ум	Характер	Красота	Здоровье	Счастье								
1	Ум												
2	62	49	50	80	50								
3	65	80	60	45	55								
4	50	65	50	50	50								
5	63	55	54	81	73								
6	63	77	51	60	73								
7	100	71	55	65	33								
8	60	78	80	69	96								
9	50	86	21	42	21								
10	40	80	47	62	48								
11	38	53	58	48	80								
12	40	55	78	79	54								
13	26	21	12	31	30								
14	75	50	50	45	50								
15	90	90	90	90	90								
16	55	66	82	67	60								

Готово

Для решения поставленной задачи методами статистического анализа Excel на панели инструментов нужно выбрать «Данные» > «Анализ данных» > «Описательная статистика» > «Корреляция»:



В диалоговом окне «Корреляция» необходимо в строке «Входной интервал» обозначить массив данных, который исследователь собирается проанализировать методом корреляции.

В нашем примере статистический пакет Excel произвел следующие расчеты:

	Столбец 1	Столбец 2	Столбец 3	Столбец 4	Столбец 5
Столбец 1	1				
Столбец 2	0,378738046	1			
Столбец 3	0,492864094	0,310419706	1		
Столбец 4	0,559649879	0,416382303	0,604736591	1	
Столбец 5	0,241840175	0,216939011	0,570103116	0,263077324	1

Исходя из данных таблицы, критические значения коэффициента корреляции Пирсона (см. приложение 2, табл. 2), статистически достоверными могут считаться зависимости в тех парах, где значение коэффициента корреляции не менее 0,396 (при $n=25$, $p<0,05$).

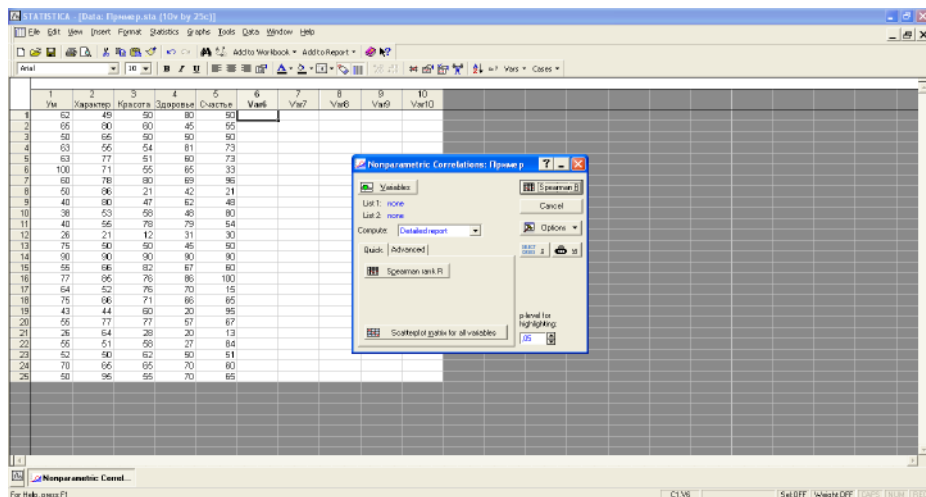
В нашем примере мы выявили пять пар переменных, между которыми образуется корреляционная связь: «1-Ум & 3-Красота», «1-Ум & 4-Здоровье», «2-Характер & 4-Здоровье», «3-Красота & 4-Здоровье», «3-Красота & 5-Счастье».

Решение в Statistica. Рассмотрим, как при помощи статистической программы Statistica провести расчет рангового коэффициента корреляции Спирмена, используемого при ненормальном распределении данных.

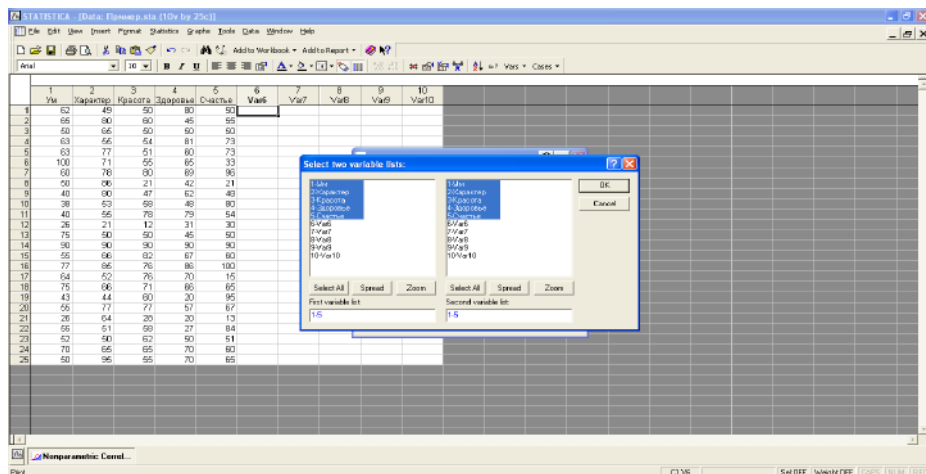
После проведения подготовительных работ с документом (см. приложение 4) он будет выглядеть следующим образом:

	1	2	3	4	5	6	7	8	9	10
	Ум	Характер	Красота	Здоровье	Счастье	Var1	Var2	Var3	Var4	Var5
1	62	49	50	80	90					
2	65	60	60	45	55					
3	58	66	50	80	60					
4	63	56	54	81	73					
5	63	77	51	60	73					
6	100	71	55	65	33					
7	60	78	80	69	96					
8	50	66	21	42	21					
9	40	80	47	62	48					
10	38	53	59	46	80					
11	40	56	78	79	54					
12	26	21	12	31	30					
13	75	50	50	45	50					
14	50	90	90	90	90					
15	55	66	62	67	60					
16	77	66	75	86	100					
17	64	52	76	70	15					
18	75	66	71	86	65					
19	43	44	60	30	95					
20	55	77	77	57	67					
21	26	64	29	20	13					
22	56	61	58	37	84					
23	53	90	62	90	51					
24	70	66	66	70	60					
25	50	96	55	70	65					

Выберите «Statistics» («Статистика») > «Nonparametrics» («Непараметрические...») > «Correlations...» («Корреляция...») и нажмите на «OK»:



В открывшемся диалоговом окне («Nonparametric Correlations») в строчке «Compute» (с англ. яз., «расчет») по умолчанию стоит «Detailed report» (с англ. яз., «детализированный отчет»). Вы можете по своему усмотрению выбрать другие режимы расчета. Нажмите кнопку «Variables» и выберите колонки с требуемыми переменными, после чего нажмите «OK»:



Далее нажмите на кнопку «Spearman R»:

STATISTICA - [Workbook11* - Spearman Rank Order Correlations (Пример)]

MD pairwise deleted
Marked correlations are significant at $p < 0.05000$

Pair of Variables	Valid N	Spearman R	t(N-2)	p-level
Ум & Ум	25	0,279043	1,393599	0,176760
Ум & Характер	25	0,281120	2,038100	0,052200
Ум & Красота	25	0,499064	2,768164	0,010940
Ум & Здоровье	25	0,247636	1,225799	0,230576
Характер & Ум	25	0,279043	1,393599	0,176760
Характер & Характер	25	0,221172	1,087641	0,288030
Характер & Красота	25	0,341238	1,741028	0,099043
Характер & Здоровье	25	0,289177	1,029071	0,319521
Красота & Ум	25	0,281120	2,038100	0,052200
Красота & Характер	25	0,221172	1,087641	0,288030
Красота & Красота	25	0,603761	2,796749	0,010246
Красота & Здоровье	25	0,603761	2,796749	0,010246
Здоровье & Ум	25	0,499064	2,768164	0,010940
Здоровье & Характер	25	0,341238	1,741028	0,099043
Здоровье & Красота	25	0,603761	2,796749	0,010246
Здоровье & Здоровье	25	0,270729	1,348738	0,190552
Счастье & Ум	25	0,247636	1,225799	0,230576
Счастье & Характер	25	0,289177	1,029071	0,319521
Счастье & Красота	25	0,603761	2,796749	0,010246
Счастье & Здоровье	25	0,270729	1,348738	0,190552
Счастье & Счастье	25	0,270729	1,348738	0,190552

В появившейся таблице в крайней левой колонке красным цветом выделены коррелируемые переменные, в колонке (Valid N) – количество участвующих в исследовании испытуемых, в колонке (Spearman R) – значение коэффициента корреляции r -Спирмена, в колонке ($t(N-2)$) – площадь пересечения распределений Стьюдента у коррелируемых переменных, в последней колонке (p-level) указан уровень значимости.

Контрольные вопросы и задания

1. Что вычисляет коэффициент корреляции? Какие возможны его значения?

2. Покажите на примере последовательность расчета коэффициента ранговой корреляции Спирмена, используя табличный редактор Excel в лицензионном пакете Microsoft Office.

3. Покажите на примере последовательность расчета коэффициента корреляции Пирсона, используя табличный редактор Excel в лицензионном пакете Microsoft Office.

4. Произведите расчет коэффициента корреляции Пирсона на тех же данных, применяя формулы и таблицы для расчета промежуточных переменных.

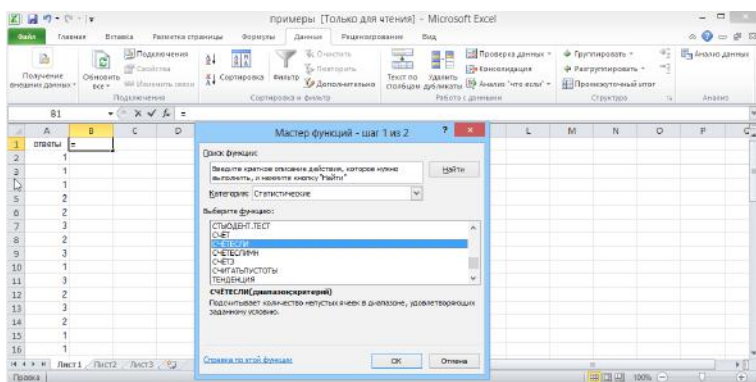
5. Покажите на примере последовательность расчета рангового коэффициента корреляции Спирмена, используя статистическую программу Statistica.

5. Вычисление частотных характеристик

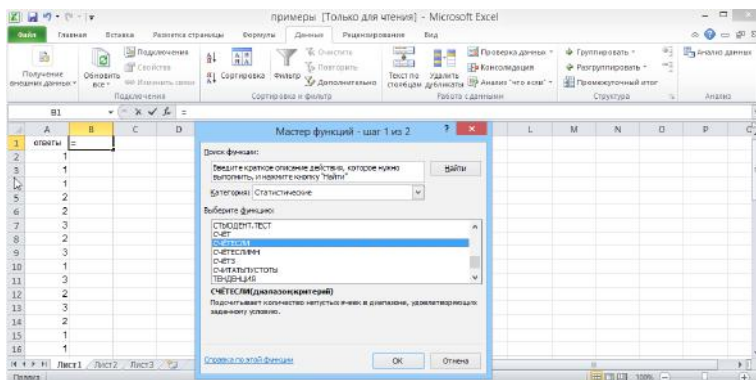
Для вычисления частотных характеристик можно при помощи «MS Office Excel» необходимо воспользоваться опцией «Вставить функцию».

Пример задачи. Исследователь должен определить количество респондентов, ответивших «1-да», «2-нет» и «3-не знаю» на следующий вопрос: «Изменился ли социально-психологический климат в служебном коллективе после проведенных в нем тренингов на сплоченность?».

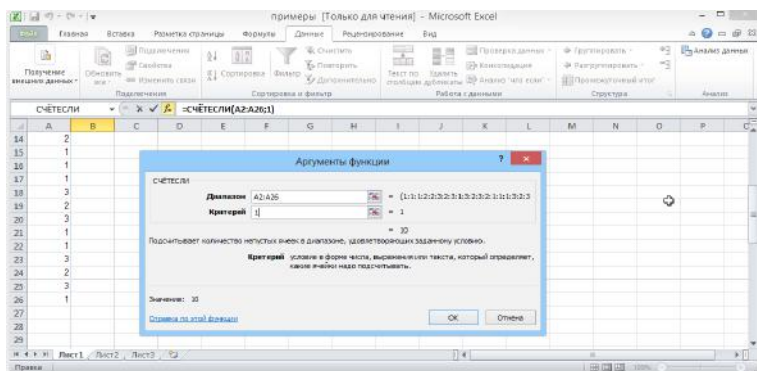
Решение. После проведения подготовительных работ с документом он будет выглядеть следующим образом:



Обязательно поставьте курсор на свободную (пустую) ячейку. Выберите опцию «Вставить функцию (f_x)», затем в диалоговом окне «Мастер функций» установите в категории «Статистические» функцию «СЧЕТЕСЛИ».



В диалоговом окне «Аргументы функции» необходимо в строке «Диапазон» обозначить массив данных, который исследователь собирается проанализировать. В строке «Критерий» обозначаем номер ответа (1, 2 или 3). В нашем примере статистический пакет Excel произвел следующие расчеты:

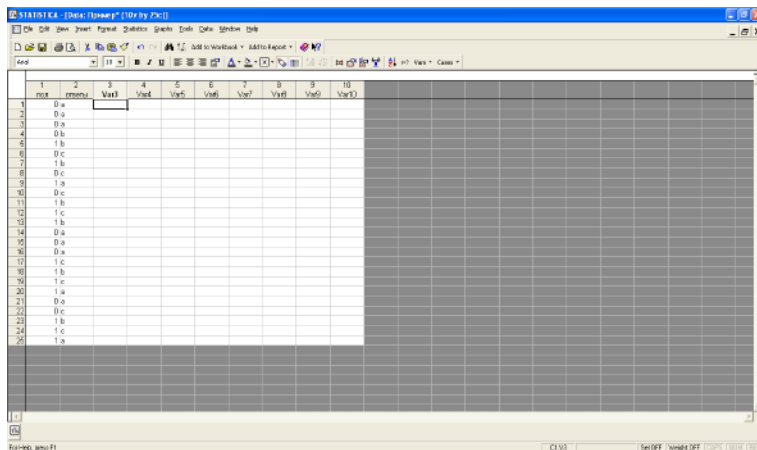


Как видно из скриншота, в нашем примере 10 человек ответили на поставленный в исследовании вопрос «да» (код 1). Таким же образом возможно узнать количество ответов на вопросы «нет» (код 2) и «не знаю» (код 3).

Для вычисления частотных характеристик при помощи программы «Статистика 6.0» необходимо воспользоваться командой «Frequency tables» в модуле «Basic statistics/Tables» («Основная статистика и таблицы»).

Пример задачи. Исследователь должен определить процентное соотношение респондентов, ответивших «да», «нет» и «не знаю» на следующий вопрос: «Считаете ли Вы уместным запретить любую рекламу табачной продукции на территории России?». Также необходимо узнать, сколько мужчин и женщин из группы респондентов ответили положительно на данный вопрос.

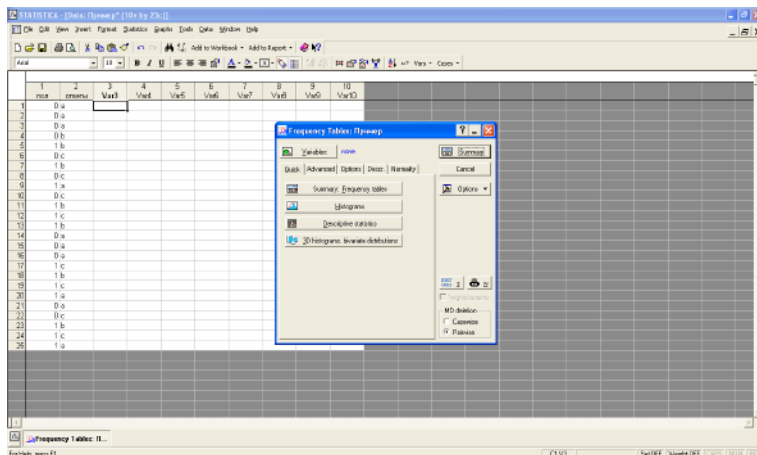
Решение. После проведения подготовительных работ с документом (см. приложение 4) он будет выглядеть следующим образом:



Необходимо отметить, что в таблице не допускается оставление пустых ячеек в столбце выбранной переменной.

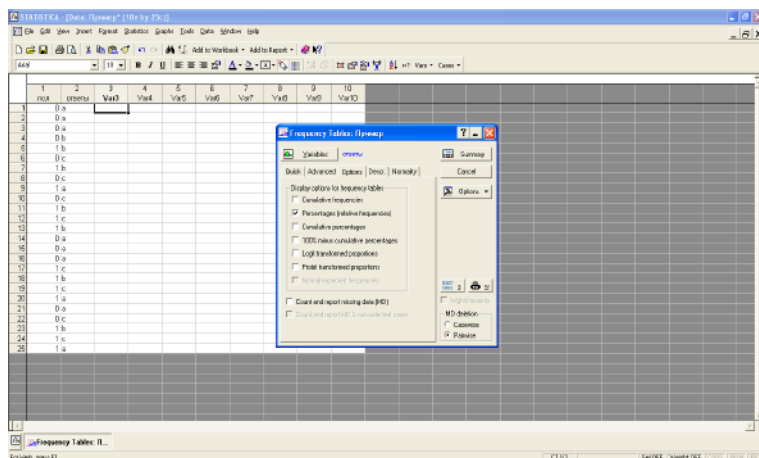
В первом столбце «0» обозначает респондента женского пола, «1» – мужского.

Выберите «Statistics» («Статистика») > «Basic statistics/Tables» («Основные статистики и таблицы») > «Frequency tables» («Таблица частот»):

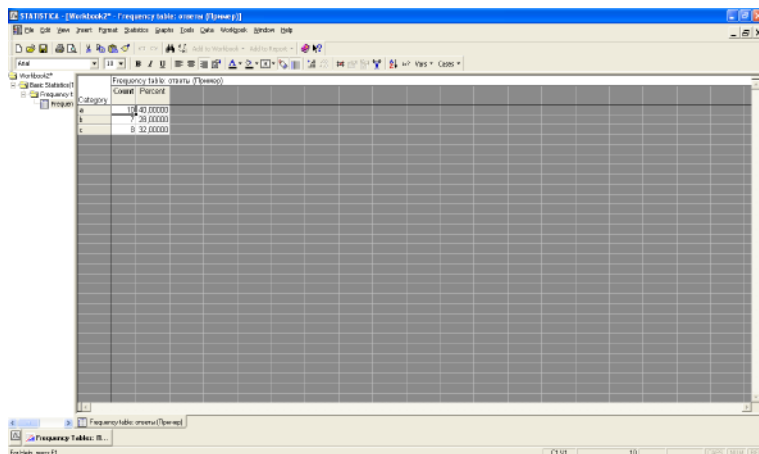


В открывшемся диалоговом окне «Frequency tables» нажмите на «Variables», выберите строчку «2-ответы», после чего нажмите «OK».

Далее нажмите на функцию «Options» («Настройки») и оставьте «галочку» только напротив «Percentages» («Проценты»):



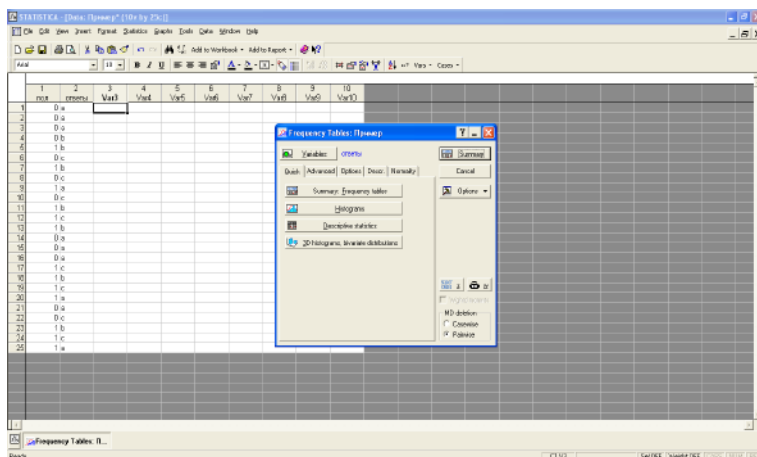
Нажмите на «Summary» («Вывод»):



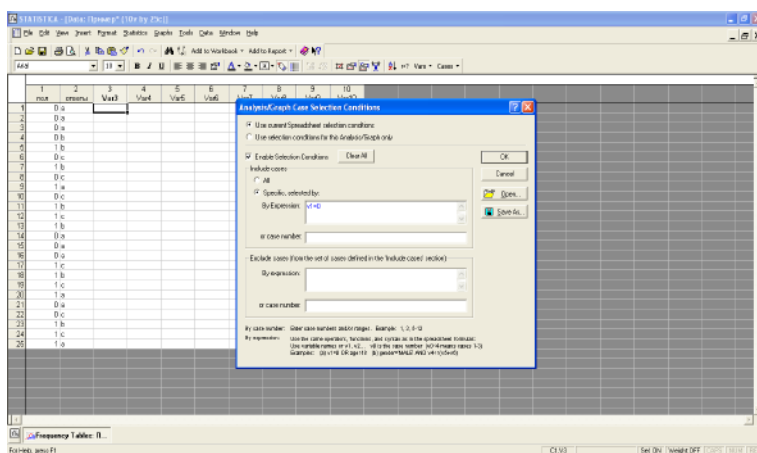
В колонке «Count» («Итого») показано число респондентов, в колонке «Percent» – их процентное соотношение. В нашем примере ответили «Да» («a») 40 % респондентов, «Нет» («b») – 28 %, «Не знаю» («c») – 32 %.

Чтобы ответить, сколько мужчин и женщин из группы респондентов ответили положительно на вопрос исследователя, необходимо в диалого-

вом окне «Frequency tables» нажать на кнопку «Select cases» («Выбор испытуемых»):



В открывшемся диалоговом окне «Analysis/Graph Case Selection Conditions» («Условия выбора испытуемых») в строчке «By Expression» («Проявление») функции «Include cases» («Выбранные испытуемые») впишите формулу $v1=0$, которая показывает, что при обработке программа выберет в столбце $v1$ респондентов только с обозначением «0», т. е. женского пола:



Далее нажмите на «ОК»:

Category	Count	Percent
f	20	30.0000
g	8	12.0000
c	4	6.0000

В нашем примере 58,3 % (семь человек) респондентов женского пола положительно ответили на вопрос исследователя.

Если в строчке «By Expression» функции «Include cases» вписать формулу $v1=1$, то в итоге программа подсчитает число и процентное соотношение ответов у мужчин:

Category	Count	Percent
a	23	23.0000
b	46	46.0000
c	4	4.0000

В нашем примере только 23 % (три человека) респондентов мужского пола положительно ответили на вопрос исследователя.

Контрольные вопросы и задания

1. Проведите расчет частотных характеристик выборочных данных, используя табличный редактор Excel в лицензионном пакете Microsoft Office.
2. Для расчета какого критерия используются частотные характеристики выборочных данных?
3. Проведите расчет критерия «хи-квадрат» Пирсона $\chi^2_{\text{эмп}}$ для выявления различий выборочных данных, предварительно рассчитав частотные характеристики сравниваемых параметров.
4. Проведите расчет частотных характеристик выборочных данных при помощи программы Statistica.

6. Регрессионный анализ

Регрессионный анализ – это один из способов (наравне с корреляционным анализом) обнаружения зависимости одного параметра от одной или нескольких независимых переменных.

Результаты регрессионного анализа позволяют предсказать, чему в среднем будет равно значение одного признака при заданном значении другого признака.

Достаточно часто связь между двумя психологическими признаками имеет линейный характер:

$$y = a + bx,$$

где y и x – анализируемые признаки;

a – свободный член уравнения; при $b = 0$ получаем $y = a$, т. е. a – это точка, в которой линия регрессии пересекается с осью OY (эту точку называют также « j -пересечением», или «*Intercept*»);

b – коэффициент регрессии, отражающий угол наклона линии регрессии. Чем больше b отличается от 0, тем сильнее связь между анализируемыми признаками.

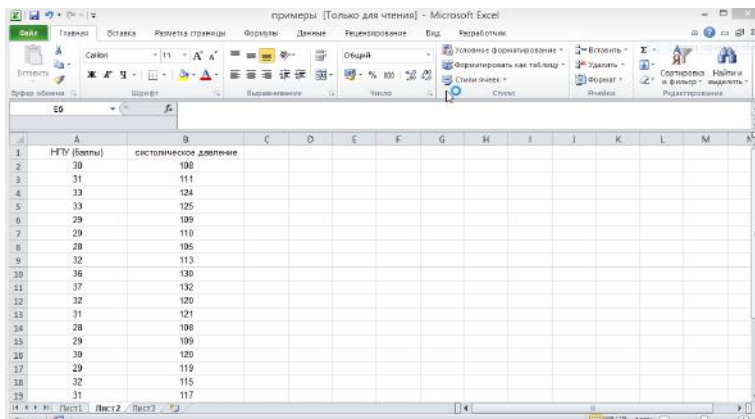
Даже если связь между психологическими признаками носит нелинейный характер (например, экспоненциальный), практически всегда можно выделить участки, хорошо аппроксимируемые линейной регрессией.

Приведенное выше уравнение можно использовать для описания связи между двумя признаками лишь при выполнении следующих обязательных условий:

- зависимость между признаками носит линейный характер;
- оба признака распределены нормально.

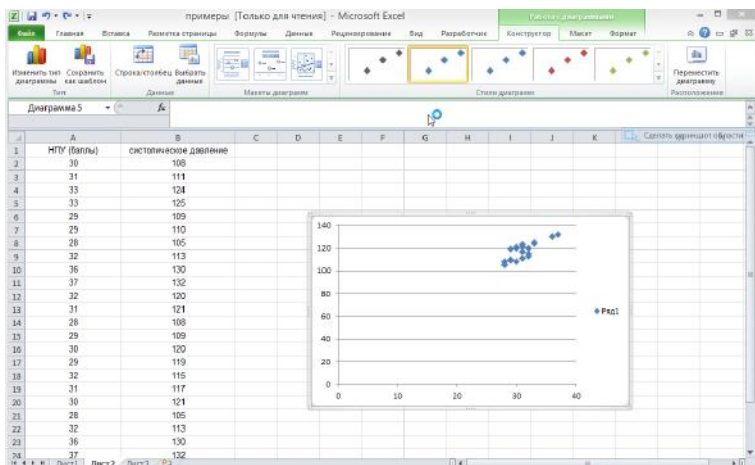
Пример задачи. Исследователь должен определить коэффициенты линейного регрессионного уравнения, описывающего связь между показателями систолического давления и нервно-психической устойчивостью (НПУ) испытуемых.

Решение в Excel. Расчет коэффициента регрессионного уравнения можно выполнить при помощи статистического пакета Excel. После проведения подготовительных работ с документом он будет выглядеть следующим образом:

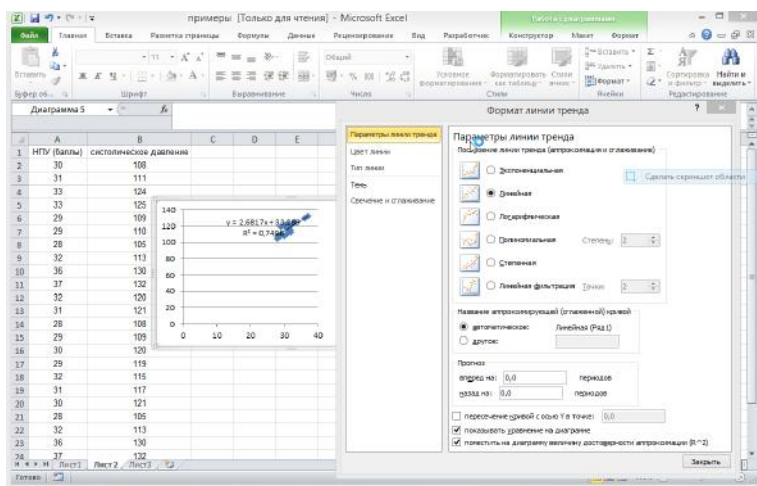


	A	B	C	D	E	F	G	H	I	J	K	L	M
1	МТВ (Ватты)	системное давление											
2	30	108											
3	31	111											
4	33	124											
5	33	125											
6	29	109											
7	29	110											
8	28	105											
9	32	113											
10	36	130											
11	37	132											
12	32	120											
13	31	121											
14	28	108											
15	29	109											
16	30	120											
17	29	119											
18	32	115											
19	31	117											

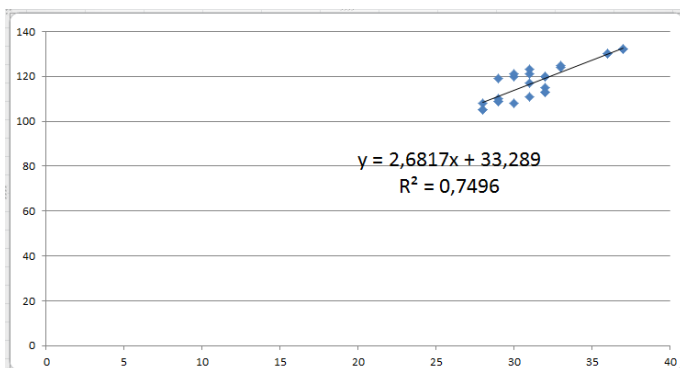
Для решения поставленной задачи методами статистического анализа Excel в первую очередь следует выделить массив данных двух переменных, а затем на панели инструментов нужно выбрать «Вставка» > «Гистограмма» > «Точечная». В итоге произведенный анализ данных позволяет построить корреляционное поле:



Щелкаем левой кнопкой мыши по любой точке на диаграмме. Потом правой. В открывшемся меню выбираем «Добавить линию тренда»:



Назначаем параметры для линии. Тип – «Линейная». Внизу необходимо поставить галочку напротив опции «Показать уравнение на диаграмме» и «Поместить на диаграмму величину достоверности аппроксимации (R^2)». Подробно рассмотрим данные проведенного регрессионного анализа.



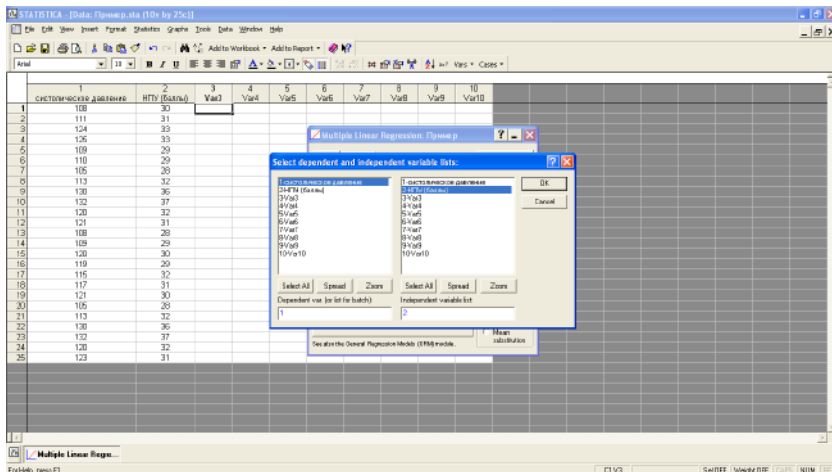
В целом построенная регрессионная модель отлично описывает связь между систолическим давлением и нервно-психической устойчивостью. R^2 (коэффициент детерминации) – очень важный показатель в регрессион-

ном анализе. Он изменяется от 0 до 1 и отражает «качество» рассчитанной регрессии, показывая долю (%) общего разброса выборочных точек, которая «объясняется» построенной регрессией (например, при $R^2 = 0,85$, следует вывод о том, что 85 % дисперсии зависимой переменной y объясняется вариацией независимой переменной x). В нашем примере коэффициент детерминации – 0,7496, или 74,96 %. Это означает, что расчетные параметры модели на 74,96 % объясняют зависимость между изучаемыми параметрами. Чем выше коэффициент детерминации, тем качественнее модель. Хороший уровень детерминации – 0,7 и выше. Плохой – меньше 0,5.

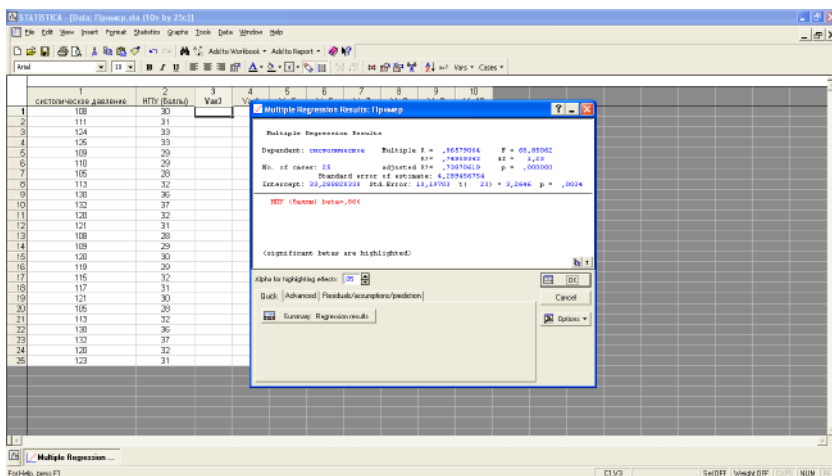
Решение в Statistica. Расчет коэффициентов регрессионных уравнений можно выполнить в нескольких модулях программы Statistica 6.0. Мы воспользуемся модулем *Multiple Regression Analysis* (Анализ множественной регрессии).

	1 систолическое давление	2 НПУ (мм/мл)	3 Var1	4 Var2	5 Var3	6 Var4	7 Var5	8 Var6	9 Var7	10 Var8	11 Var9	12 Var10
1	108	31										
2	111	31										
3	124	33										
4	126	33										
5	109	29										
6	110	29										
7	106	28										
8	113	32										
9	130	36										
10	132	37										
11	120	32										
12	121	31										
13	108	28										
14	109	29										
15	120	30										
16	119	29										
17	116	32										
18	117	31										
19	121	30										
20	106	28										
21	113	32										
22	130	36										
23	132	37										
24	120	32										
25	123	31										

В появившемся окне нажмите на кнопку *Variables* и укажите, какая из анализируемых переменных является зависимой (*Dependent variable*), а какая – независимой (*Independent variable*) (в нашем примере систолическое давление зависит от НПУ).



Нажмите на кнопку «ОК».



В итоге появится окно, которое уже на данном этапе анализа содержит некоторые важные его результаты:

- а) *Dependent*: имя зависимой переменной;
- б) *No. of cases*: число наблюдений;
- в) *Intercept*: значение свободного члена регрессионного уравнения;

г) *Std. error*: стандартная ошибка свободного члена регрессионного уравнения;

д) *Multiple R*: коэффициент множественной корреляции;

е) *R*: коэффициент детерминации. Это очень важный показатель в регрессионном анализе. Он изменяется от 0 до 1 и отражает «качество» рассчитанной регрессии, показывая долю (%) общего разброса выборочных точек, которая «объясняется» построенной регрессией (например, при $R^2 = 0,85$, следует вывод о том, что 85 % дисперсии зависимой переменной y объясняется вариацией независимой переменной x);

ж) *Adjusted R²*: скорректированный на число степеней свободы коэффициент детерминации ($Adjusted\ R-square = 1 - (1 - R-square) \times [n/(n - p)]$), где n – число наблюдений, p – число независимых переменных плюс 1);

з) *Standard error of estimate*: параметр, отражающий степень разброса выборочных значений относительно линии регрессии;

и) *F*, *dfup*: F-критерий, число степеней свободы, принятое при его расчете, и вероятность ошибки для нулевой гипотезы F-теста. F-тест в регрессионном анализе применяется для оценки статистической значимости модели. При $p < 0,05$ можно заключить, что рассчитанная регрессия удовлетворительно описывает связь между исследуемыми признаками;

к) *t(df)* и *p*: критерий Стьюдента t используется для проверки нулевой гипотезы о равенстве 0 свободного члена регрессионного уравнения. P – вероятность ошибки для этой нулевой гипотезы;

л) *beta*: стандартизованный коэффициент регрессии – это коэффициент регрессии, который мы получили бы в случае предварительной стандартизации обеих переменных (т. е. при таком преобразовании, когда их средние значения стали бы равны 0, а стандартные отклонения – 1). Расчет *beta* позволяет оценить, в какой степени значения зависимой переменной определяются значениями независимой переменной. *Beta* может оказаться особенно полезным показателем при включении в анализ нескольких независимых переменных, выражающихся в разных единицах измере-

ния – в таком случае коэффициент отражал бы удельный вклад каждой из этих переменных в вариацию зависимой переменной. При наличии одной независимой переменной коэффициент β идентичен $Multiple\ R$.

Нажмите кнопку «Summary: Regression results» (Результаты регрессионного анализа). Появится таблица с результатами анализа, в которой:

- а) $Beta$: стандартизованный коэффициент регрессии;
- б) $Std.\ err.\ of\ beta$: стандартная ошибка стандартизованного коэффициента регрессии;
- в) B : один из самых важных столбцов в этой таблице, поскольку именно он содержит искомые значения свободного члена регрессионного уравнения (в строке *Intercept*) и коэффициента регрессии (нижняя строка таблицы);
- г) $Std.\ err.\ of\ B$: стандартные ошибки коэффициентов уравнения;
- д) $t(df)$: значения t-критерия Стьюдента, который используется для проверки гипотезы о равенстве обоих коэффициентов уравнения 0;
- е) $p\text{-level}$: вероятность ошибки для нулевой гипотезы о равенстве коэффициентов уравнения нулю.

	Beta	Std. Err. of Beta	B	Std. Err. of B	t(df)	p-level
Intercept			33.26881	10.10701	3.28662	0.000489
HTV (Battu)	0.807179	0.104342	2.66189	0.32316	8.250626	0.000000

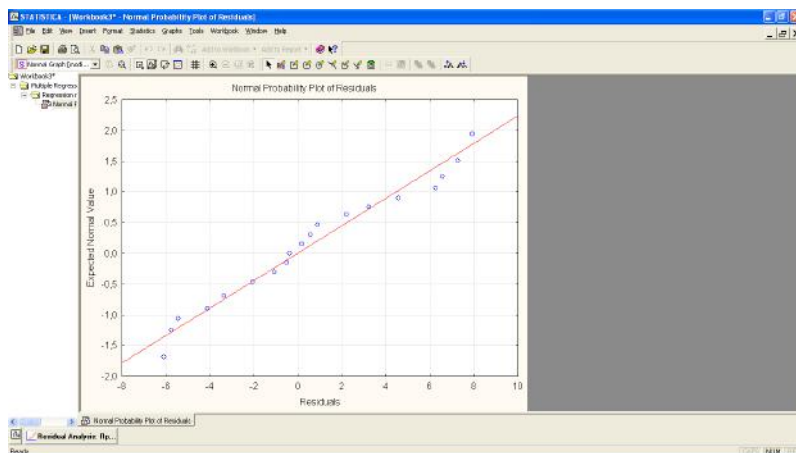
Из итоговой таблица видно, что оба коэффициента регрессии статистически значимо отличаются от 0 ($p < 0,05$) и что в целом построенная

регрессионная модель отлично описывает связь между систолическим давлением и нервно-психической устойчивостью. Само же рассчитанное уравнение мы можем записать следующим образом:

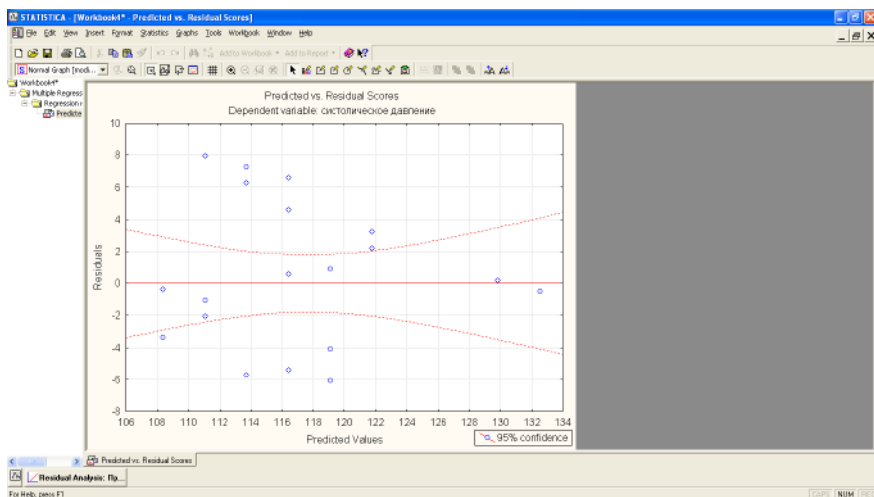
$H = 2,682 \times A + 33,289$, где H - давление, A – нервно-психическая устойчивость человека.

Важной частью регрессионного анализа является *анализ остатков* (остатки представляют собой разности между наблюдаемыми значениями зависимой переменной и теми ее значениями, которые предсказываются регрессионной моделью). Он запускается путем нажатия кнопки «Perform residual analysis» (Выполнить анализ остатков) на закладке «Residuals / Assumptions / Predictions» (Остатки / Условия / Предсказания).

Первое, что нужно проверить в отношении остатков, – это нормальность их распределения. Для этого на закладке «Quick» подмодуля анализа остатков нажмите кнопку «Normal plot of residuals», чтобы построить график нормальных вероятностей. Если точки на этом графике достаточно тесно укладываются вдоль теоретически ожидаемой прямой, то можно заключить, что остатки распределяются нормально. Иначе линейная регрессионная модель для анализируемых переменных будет неприменима.



Второе условие в отношении остатков состоит в том, что их дисперсия должна оставаться неизменной во всем диапазоне значений анализируемых переменных. Для проверки этого условия на закладке «Scatterplots» (Диаграммы рассеяния) нажмите кнопку «Predicted vs. Residuals», чтобы построить график зависимости значений остатков от предсказываемых моделью значений зависимой переменной. Если проверяемое условие выполняется, то точки на этом графике будут располагаться хаотично, не проявляя никакой закономерности. Если же в расположении точек имеется тенденция (разброс увеличивается слева направо, точки тесно укладываются вдоль прямой и т. п.), линейный регрессионный анализ также неприменим.



В рассмотренном примере оба условия в отношении остатков выполняются, что еще раз подтверждает адекватность рассчитанной регрессионной модели для описания связи между артериальным давлением и нервно-психической устойчивостью.

Контрольные вопросы и задания

1. Какую задачу решает регрессионный анализ при обработке эмпирических данных исследования?
2. Проведите расчет регрессионного уравнения по выборочным данным двух переменных, полученных в ходе эмпирического исследования, применяя табличный редактор Excel в лицензионном пакете Microsoft Office и программу Statistica.
3. Какую роль играет коэффициент детерминации R^2 в регрессионном анализе?

7. Факторный анализ

Факторный анализ – статистический метод, который используется при обработке больших массивов экспериментальных данных. Задачами факторного анализа являются: сокращение числа переменных (редукция данных) и определение структуры взаимосвязей между переменными, т. е. классификация переменных, поэтому факторный анализ используется как метод сокращения данных или как метод структурной классификации [2].

Важное отличие факторного анализа от описанных выше методов заключается в том, что его нельзя применять для обработки первичных, или, как говорят, «сырых», экспериментальных данных, т. е. полученных непосредственно при обследовании испытуемых. Материалом для факторного анализа служат корреляционные связи, а точнее – коэффициенты корреляции Пирсона, которые вычисляются между переменными (т. е. психологическими признаками), включенными в обследование.

Главное понятие факторного анализа – фактор. Это искусственный статистический показатель, возникающий в результате специальных преобразований таблицы коэффициентов корреляции между изучаемыми психологическими признаками, или матрицы интеркорреляций. Процедура извлечения факторов из матрицы интеркорреляций называется факторизацией матрицы. В результате факторизации из корреляционной матрицы может быть извлечено разное количество факторов вплоть до числа, равного количеству исходных переменных. Однако факторы, выделяемые в результате факторизации, как правило, неравноценны по своему значению [1, 3].

Элементы факторной матрицы называются «факторными нагрузками, или весами»; и они представляют собой коэффициенты корреляции данного фактора со всеми показателями, использованными в исследовании.

Факторная матрица показывает, какие переменные образуют каждый фактор. Это связано прежде всего с уровнем значимости факторного веса. По традиции минимальный уровень значимости коэффициентов корреляции в факторном анализе берется равным 0,4 или даже 0,3 (по абсолютной величине), поскольку нет специальных таблиц, по которым можно было бы определить критические значения для уровня значимости в факторной матрице. Следовательно, самый простой способ увидеть, какие переменные «принадлежат» фактору – это значит отметить те из них, которые имеют нагрузки выше, чем 0,4 (или меньше, чем 0,4). В некоторых компьютерных пакетах иногда уровень значимости факторного веса определяется самой программой и устанавливается на более высоком уровне, например 0,7.

Перед факторизацией необходимо провести процедуру вращения факторов, которая изменяет положение факторов по отношению к переменным таким образом, что получаемое решение легко интерпретировать. Цель вращения – преобразовать факторную матрицу таким образом, чтобы получилась простая структура, в которой каждый фактор имеет некоторое количество больших нагрузок и некоторое малых, и подобно этому каждая переменная имеет существенные нагрузки только по некоторым факторам.

Переменные могут иметь разную степень общности с факторами. Переменная с большей общностью имеет значительную степень перекрытия (большую долю дисперсии) с одним или несколькими факторами. Низкая общность подразумевает, что все корреляции между переменными и факторами невелики. Это означает, что ни один из факторов не имеет совпадающей доли вариативности с данной переменной. Низкая общность может свидетельствовать о том, что переменная измеряет нечто качественно отличающееся от других переменных, включенных в анализ. Например, одна переменная, связанная с оценкой мотивации

среди заданий, оценивающих способности, будет иметь общность с факторами способностей, близкую к нулю.

С помощью такой характеристики, как собственное значение фактора, можно определить относительную значимость каждого из выделенных факторов. Для этого надо вычислить, какую часть дисперсии (вариативности) объясняет каждый фактор [2].

Как показывает практика, психологи предпочитают использовать метод варимакс [2]. При использовании данного метода минимизируется количество переменных, имеющих высокие нагрузки на данный фактор, при этом максимально увеличивается дисперсия фактора. Это способствует упрощению описания фактора за счет группировки вокруг него только тех переменных, которые в большей степени связаны с ним, чем остальные.

Разработано несколько приемов для выбора «правильного» числа факторов из корреляционной матрицы. Определение числа выделяемых факторов, вероятно, наиболее важное решение, которое необходимо принять при проведении факторного анализа. Неверное решение может привести к бессмысленным результатам при обработке самого четкого набора данных. Нет ничего страшного в том, чтобы попытаться выполнить несколько вариантов анализа, базирующегося на разном числе факторов, и использовать несколько различных приемов, определяющих выбор факторов.

Первые руководящие принципы – это теория, здравый смысл, а также прошлый опыт. При этом психолог должен установить:

- не способствует ли увеличение числа факторов уменьшению доли нагрузок в диапазоне от -0,4 до +0,4? Если это так, то это увеличение, скорее всего, не имеет смысла;

- не появляются ли какие-либо большие корреляции между факторами при осуществлении вращений;

– не разделились ли какие-либо хорошо известные факторы на две или большее количество частей.

Факторный анализ может быть уместен, если выполняются следующие условия.

1. Нельзя факторизовать качественные данные, полученные по шкале наименований, например, такие, как цвет волос (черный / каштановый / рыжий) и т. п.

2. Все переменные должны быть независимыми, а их распределение должно приближаться к нормальному.

3. Связи между переменными должны быть приблизительно линейны или, по крайней мере, не иметь явно криволинейного характера.

4. В исходной корреляционной матрице должно быть несколько корреляций по модулю выше 0,3. В противном случае достаточно трудно извлечь из матрицы какие-либо факторы.

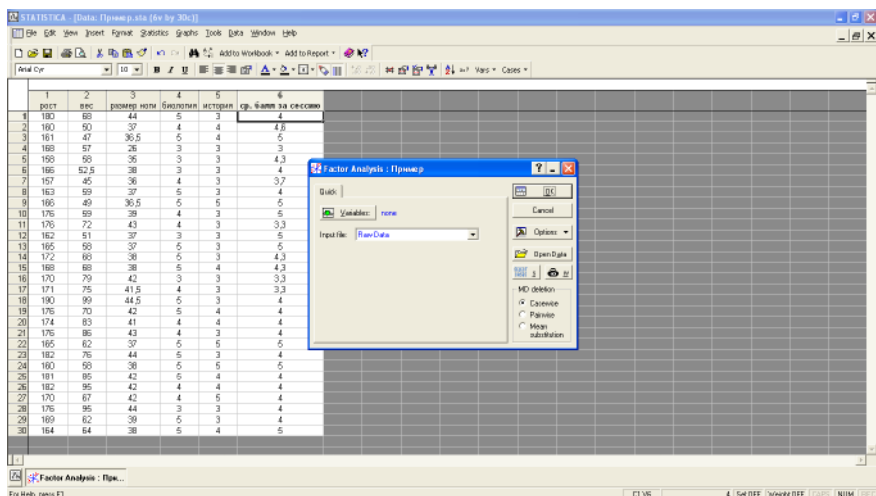
5. Выборка испытуемых должна быть достаточно большой. Рекомендации экспертов варьируют. Наиболее жесткая точка зрения рекомендует не применять факторный анализ, если число испытуемых меньше 100, поскольку стандартные ошибки корреляции в этом случае окажутся слишком велики. Однако если факторы хорошо определены (например, с нагрузками 0,7, а не 0,3), экспериментатору нужна меньшая выборка, чтобы выделить их. Кроме того, если известно, что полученные данные отличаются высокой надежностью (например, используются валидные тесты), то можно анализировать данные и по меньшему числу испытуемых.

Пример задачи. Необходимо выделить два фактора из шести имеющихся переменных: рост, вес, размер ноги курсанта, его отметки за вступительный экзамен по биологии и истории, а также средний балл за сессию.

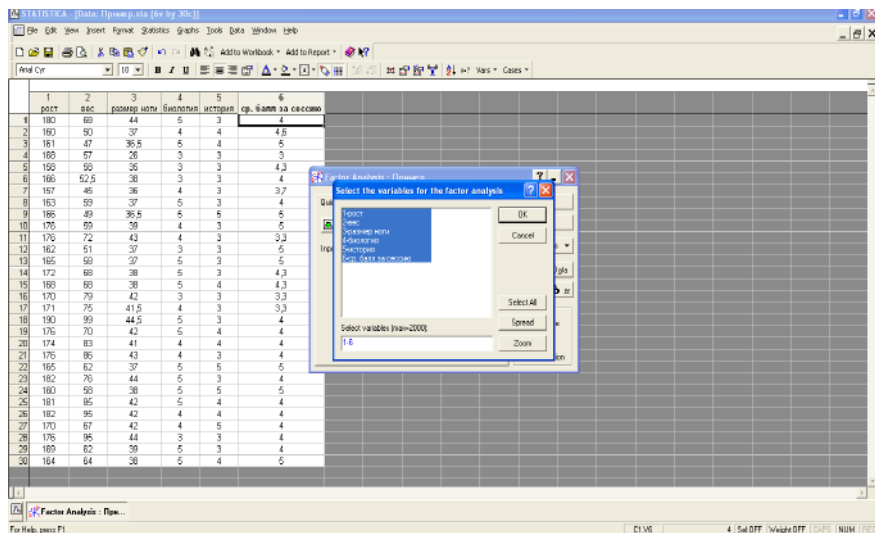
Решение. После проведения подготовительных работ с документом (см. приложение 4) он будет выглядеть следующим образом:

	1. полет	2. вес	3. размер ноги	4. биологический	5. исторический	6. ср. баллы за сессию
1	180	69	44	5	3	4
2	180	90	37	4	4	4.8
3	181	47	36.5	5	4	5
4	180	57	26	3	3	3
5	190	99	35	3	3	4.3
6	185	52.5	38	3	3	4
7	157	45	36	4	3	3.7
8	163	59	37	5	3	4
9	195	49	36.5	5	5	5
10	176	69	39	4	3	5
11	176	72	43	4	3	3.3
12	182	51	37	3	3	5
13	165	99	37	5	3	5
14	172	69	38	5	3	4.3
15	180	80	38	5	4	4.3
16	170	79	42	3	3	3.3
17	171	75	41.5	4	3	3.3
18	190	99	44.5	5	3	4
19	176	70	42	5	4	4
20	174	83	41	4	4	4
21	176	85	43	4	3	4
22	165	62	37	5	5	5
23	182	76	44	5	3	4
24	180	99	38	5	5	5
25	181	85	42	5	4	4
26	182	95	42	4	4	4
27	170	67	42	4	5	4
28	176	95	44	3	3	4
29	169	62	39	5	3	4
30	164	64	38	5	4	5

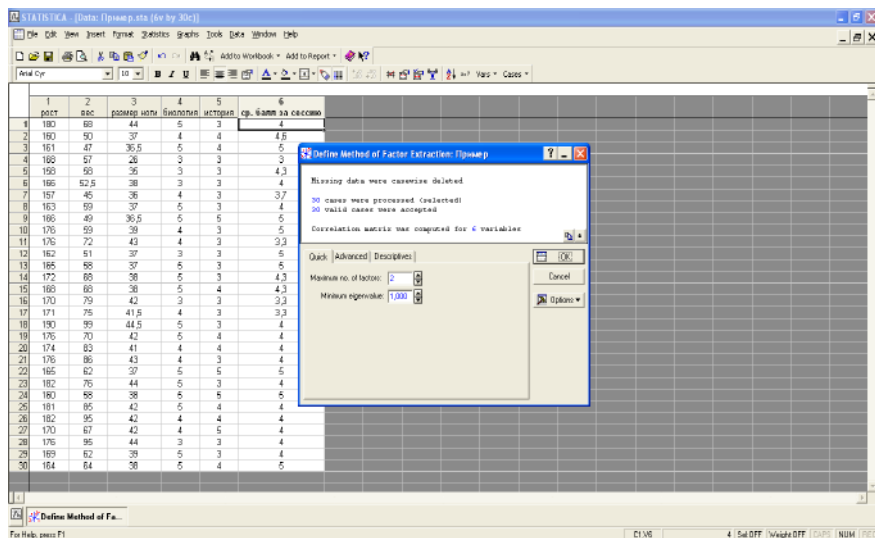
Выберите «Statistics» («Статистика») > «Multivariate Exploratory Techniques» («Многомерные разведочные технологии») > «Factor Analysis» («Факторный анализ»):



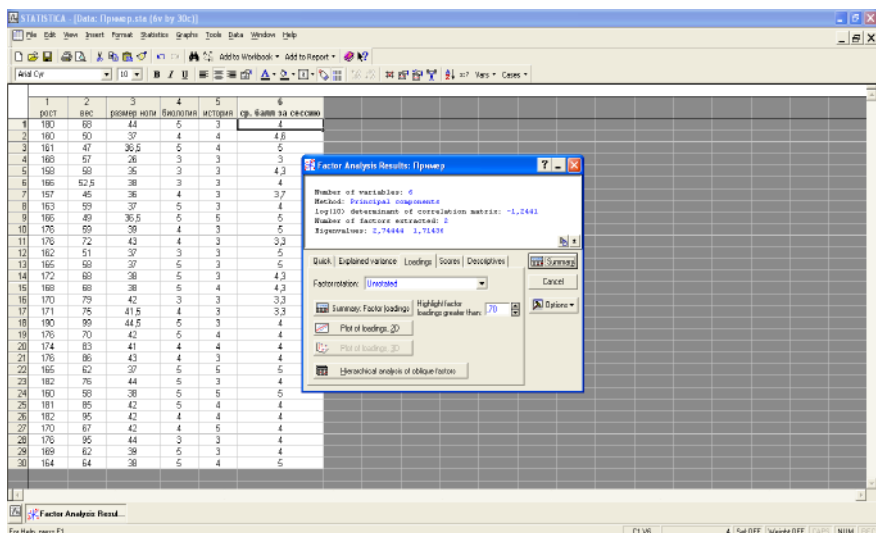
В открывшемся диалоговом окне «Factor Analysis» нажмите кнопку «Variables» и выберите колонки с требуемыми переменными, в нашем примере это шесть позиций:



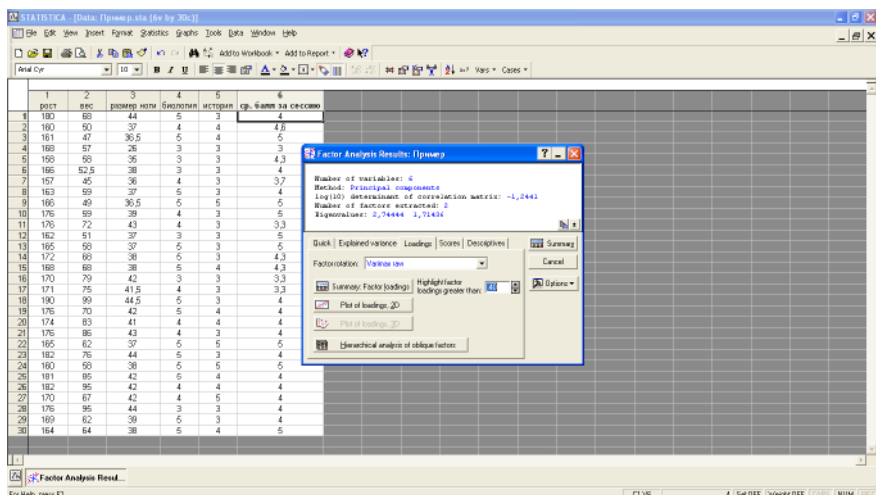
После чего нажмите «OK» и еще раз «OK»:



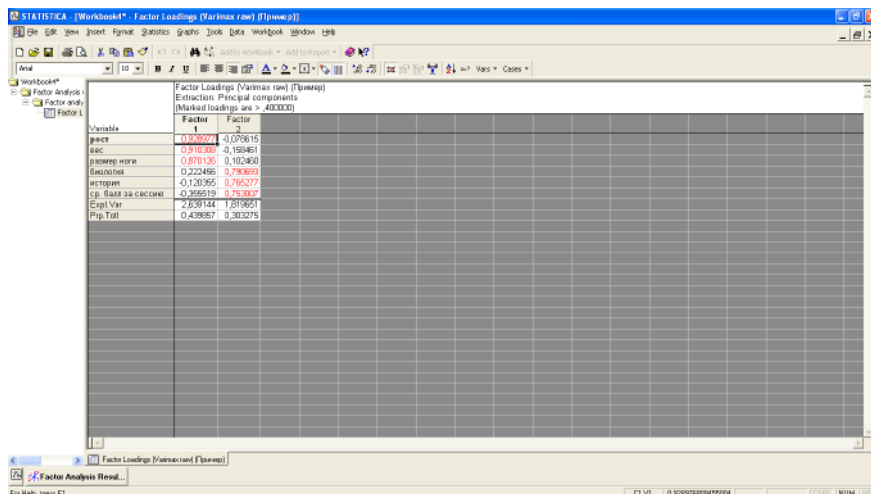
По умолчанию в строчке «Maximum no. of factors» («Максимальное количество факторов») установлено два выделяемых фактора. Нажимаете «OK»:



В строчке «Factor rotation» («Вращение фактора») установите «Varimax raw» («Варимак сырых данных»). В строчке «Highlight factor loading greater than» («Факторная нагрузка больше, чем») установите минимально допустимое значение – 0,4:



Нажмите на «Summary»:



В нашем примере мы выявили два очевидных фактора, которые условно можно назвать «Конституция (телосложение)» и «Успеваемость». В фактор «Конституция» вошли такие переменные, как «рост», «вес» и «размер ноги». В фактор «Успеваемость» – «биология», «история» и «средний балл за сессию». Стоит отметить, что у всех переменных факторная нагрузка оказалась выше 0,7, что свидетельствует о их весомой статистической принадлежности соответствующему фактору. В нашем примере фактор I объясняет около 44 % информации (дисперсии) в исходной корреляционной матрице, фактор II – 30 %. В сумме это будет составлять 74 %. Таким образом, два общих фактора, будучи объединены, объясняют только 74 % дисперсии показателей исходной корреляционной матрицы. В результате факторизации оставшаяся часть информации в исходной корреляционной матрице (26 %) была принесена в жертву построению двухфакторной модели.

Контрольные вопросы и задания

1. Какие задачи решает факторный анализ при обработке эмпирических данных исследования?
2. Проведите факторный анализ по выборочным данным не менее шести переменных, полученных в ходе эмпирического исследования, применяя программу Statistica.
3. Какую роль играет факторная нагрузка при проведении факторного анализа? Для чего нужна процедура вращения факторов?

БИБЛИОГРАФИЧЕСКИЙ СПИСОК

1. Дружинин В. Н. Экспериментальная психология. – СПб. : Питер, 2011. – 320 с.
2. Ермолаев-Томин О. Ю. Математические методы в психологии : учебник для бакалавров. – М. : Юрайт, 2013. – 511 с.
3. Корнилова Т. В. Экспериментальная психология : в 2 ч. Ч.1. – 3-е изд., перераб. и доп. – М. : Юрайт, 2016. – 383 с.
4. Куликов Л. В. Психологическое исследование: методические рекомендации по проведению. – СПб. : ООО «Речь», 2001. – С. 90–92.
5. Мاستицкий С. Э. Методическое пособие по использованию программы Statistica при обработке данных биологических. – Мн. : РУП «Институт рыбного хозяйства», 2009. – 76 с.
6. Носс И. Н. Экспериментальная психология: учебник и практикум для академического бакалавриата. – М. : Юрайт, 2015. – 321 с.
7. Практикум по общей, экспериментальной и прикладной психологии / под ред. А. А. Крылова, С. А. Маничева. – 2-е изд., доп. и перераб. – СПб., 2003. – 560 с.
8. Сидоренко Е. В. Методы математической обработки в психологии. – СПб. : ООО «Речь», 2000. – 350 с.
9. Субботина А. В., Гржибовский А. М. Описательная статистика и проверка нормальности распределения количественных данных // Экология человека. – 2014. – № 2. – С. 51–57.
10. Халифян А. А. Statistica 6. Статистический анализ данных : учебник. – 3-е изд. – М. : Бином-Пресс, 2008. – 512 с.
- Худяков А. И. Экспериментальная психология в схемах и комментариях. – СПб. : Питер, 2008. – 320 с.
11. Волков Б. С., Волкова Н. В., Губанов А. В. Методология и методы психологического исследования : учебное пособие для вузов. – М. : Академический проект, 2010. – 382 с.

ПРИЛОЖЕНИЯ

Приложение 1

Таблица 1

**Значения критерия t_{st} для отбраковки выпадающих вариантов
при разных уровнях значимости (p)**

	p				p		
n	0,05	0,01	0,001	n	0,05	0,01	0,001
-	-	-	-	-	-	-	-
5	3,04	5,04	9,43	21	2,145	2,932	3,979
6	2,78	4,36	7,41	25	2,105	2,852	3,819
7	2,62	3,96	6,37	30	2,079	2,802	3,719
8	2,51	3,71	5,73	35	2,061	2,768	3,652
9	2,43	3,54	5,31	40	2,048	2,742	3,602
10	2,37	3,41	5,01	45	2,038	2,722	3,565
11	2,33	3,31	4,79	50	2,030	2,707	3,532
12	2,29	3,23	4,62	60	2,018	2,683	3,492
13	2,26	3,17	4,48	70	2,009	2,667	3,462
14	2,24	3,12	4,37	80	2,003	2,655	3,439
15	2,22	3,08	4,28	90	1,998	2,646	3,423
16	2,20	3,04	4,20	100	1,994	2,639	3,409
17	2,18	3,01	4,13	∞	1,960	2,576	3,291
18	2,17	2,98	4,07				

Таблица 2

**Значение функции $f(o_i)$
(ординаты нормальной кривой)**

O_i	Сотые доли O_i									
	0	1	2	3	4	5	6	7	8	9
0,0	3989	3989	3989	3988	3986	3984	3982	3980	3977	3973
0,1	3970	3965	3961	3956	3951	3945	3939	3932	3925	3918
0,2	3910	3902	3894	3885	3876	3867	3857	3847	3836	3825
0,3	3814	3802	3790	3778	3765	3752	3739	3726	3712	3697
0,4	3683	3668	3653	3637	3621	3605	3589	3572	3555	3538
0,5	3521	3503	3485	3467	3448	3429	3410	3391	3372	3352
0,6	3332	3312	3292	3271	3251	3230	3209	3187	3166	3144
0,7	3123	3101	3079	3056	3034	3011	2989	2966	2943	2920
0,8	2897	2874	2850	2827	2803	2780	2756	2732	2709	2685
0,9	2661	2637	2613	2589	2565	2541	2516	2492	2468	2444
1,0	2420	2396	2371	2347	2323	2299	2275	2251	2227	2203
1,1	2179	2155	2131	2107	2083	2059	2036	2012	1989	1965
1,2	1942	1919	1895	1872	1849	1826	1804	1781	1758	1736
1,3	1714	1691	1669	1647	1626	1604	1582	1561	1539	1518
1,4	1497	1476	1456	1435	1415	1394	1374	1354	1334	1315
1,5	1295	1276	1257	1238	1219	1200	1182	1163	1145	1127
1,6	1109	1092	1074	1057	1040	1023	1006	0989	0973	0957
1,7	0940	0925	0909	0893	0878	0863	0848	0833	0818	0804
1,8	0790	0775	0761	0748	0734	0721	0707	0694	0681	0669
1,9	0656	0614	0632	0620	0608	0596	0584	0573	0562	0551
2,0	0540	0529	0519	0508	0498	0488	0478	0468	0459	0449

O_i	Сотые доли O_i									
	0	1	2	3	4	5	6	7	8	9
2,2	0355	0347	0339	0332	0325	0317	0310	0303	0297	0290
2,3	0283	0277	0270	0264	0258	0252	0246	0241	0235	0229
2,4	0224	0219	0213	0208	0203	0198	0194	0189	0184	0180
2,5	0175	0171	0167	0163	0158	0154	0151	0147	0143	0139
2,6	0136	0132	0129	0126	0122	0119	0116	0113	0110	0107
2,7	0104	0101	0099	0096	0093	0091	0088	0086	0084	0081
2,8	0079	0077	0075	0073	0071	0069	0067	0065	0063	0061
2,9	0060	0058	0056	0055	0053	0051	0050	0048	0047	0046
3,0	0044	0043	0042	0040	0039	0038	0037	0036	0035	0034
3,1	0033	0032	0031	0030	0029	0028	0027	0026	0025	0025
3,2	0024	0023	0022	0022	0021	0020	0020	0019	0018	0018
3,3	0017	0017	0016	0016	0015	0015	0014	0014	0013	0013
3,4	0012	0012	0012	0011	0011	0010	0010	0010	0009	0009
3,5	0009	0008	0008	0008	0008	0007	0007	0007	0007	0006
3,6	0006	0006	0006	0005	0005	0005	0005	0005	0005	0004
3,7	0004	0004	0004	0004	0004	0004	0003	0003	0003	0003
3,8	0003	0003	0003	0003	0003	0002	0002	0002	0002	0002
3,9	0002	0002	0002	0002	0002	0002	0002	0002	0001	0001
4,0	0001	0001	0001	0000	0000	0000	0000	0000	0000	0000

Примечание: значения вероятности p даны числами после запятой.

Таблица 3

Критические значения критерия χ^2 для уровней статистической значимости $p \leq 0,05$ и $p \leq 0,01$ при разном числе степеней свободы V

Различия между двумя распределениями могут считаться достоверными, если $\chi^2_{\text{эмп}}$ достигает или превышает $\chi^2_{0,05}$, и тем более достоверными, если $\chi^2_{\text{эмп}}$ достигает или превышает $\chi^2_{0,01}$
(по Л. Н. Большеву, Н. В. Смирнову, 1983).

p			p			p		
V	0,05	0,01	V	0,05	0,01	V	0,05	0,01
1	3,841	6,635	35	49,802	57,342	69	89,391	99,227
2	5,991	9,210	36	50,998	58,619	70	90,631	100,425
3	7,815	11,345	37	52,192	59,892	71	91,670	101,621
4	9,488	13,277	38	53,384	61,162	72	92,808	102,816
5	11,070	15,086	39	54,572	62,428	73	93,945	104,010
6	12,592	16,812	40	55,758	63,691	74	95,081	105,202
7	14,067	18,475	41	56,942	64,950	75	96,217	106,393
8	15,507	20,090	42	58,124	66,206	76	97,351	107,582
9	16,919	21,666	43	59,304	67,459	77	98,484	108,771
10	18,307	23,209	44	60,481	68,709	78	99,617	109,958
11	19,675	24,725	45	61,656	69,957	79	100,749	111,144
12	21,026	26,217	46	62,830	71,201	80	101,879	112,329
13	22,362	27,688	47	64,001	72,443	81	103,010	113,512
14	23,685	29,141	48	65,171	73,683	82	104,139	114,695
15	24,996	30,578	49	66,339	74,919	83	105,267	115,876
16	26,296	32,000	50	67,505	76,154	84	106,395	117,057
17	27,587	33,409	51	68,669	77,386	85	107,522	118,236

Продолжение табл. 3

<i>p</i>			<i>p</i>			<i>p</i>		
<i>V</i>	0,05	0,01	<i>V</i>	0,05	0,01	<i>V</i>	0,05	0,01
18	28,869	34,805	52	69,832	78,616	86	108,648	119,414
19	30,144	36,191	53	70,993	79,843	87	109,773	120,591
20	31,410	37,566	54	72,153	81,069	88	110,898	121,767
21	32,671	38,932	55	73,311	82,292	89	112,022	122,942
22	33,924	40,289	56	74,468	83,513	90	113,145	124,116
23	35,172	41,638	57	75,624	84,733	91	114,268	125,289
24	36,415	42,980	58	76,778	85,950	92	115,390	126,462
25	37,652	44,314	59	77,931	87,166	93	116,511	127,633
26	38,885	45,642	60	79,082	88,379	94	117,632	128,803
27	40,113	46,963	61	80,232	89,591	95	118,752	129,973
28	41,337	48,278	62	81,381	90,802	96	119,871	131,141
29	42,557	49,588	63	82,529	92,010	97	120,990	132,309
30	43,773	50,892	64	83,675	93,217	98	122,108	133,476
31	44,985	52,191	65	84,821	94,422	99	123,225	134,642
32	46,194	53,486	66	85,965	95,626	100	124,342	135,807
33	47,400	54,776	67	87,108	96,828			
34	48,602	56,061	68	88,250	98,028			

Таблица 4

**Значения критерия t Стьюдента при различных уровнях
значимости (p)**

Число степеней свободы d	Уровень значимости		
	0,05	0,01	0,001
1	12,71	63,66	-
2	4,30	9,93	31,60
3	3,18	5,84	12,94
4	2,78	4,60	8,61
5	2,57	4,03	6,86
6	2,45	3,71	5,96
7	2,37	3,50	5,41
8	2,31	3,36	5,04
9	2,26	3,25	4,78
10	2,23	3,17	4,59
11	2,20	3,11	4,44
12	2,18	3,06	4,32
13	2,16	3,01	4,22
14	2,15	2,98	4,14
15	2,13	2,95	4,07
16	2,12	2,92	4,02
17	2,11	2,90	3,97
18	2,10	2,88	3,92
19	2,09	2,86	3,88

Продолжение табл. 4

Число степеней свободы d	Уровень значимости		
	0,05	0,01	0,001
20	2,09	2,85	3,85
21	2,08	2,83	3,82
22	2,07	2,82	3,79
23	2,07	2,81	3,77
24	2,06	2,80	3,75
25	2,06	2,79	3,73
26	2,06	2,78	3,71
27	2,05	2,77	3,69
28	2,05	2,76	3,67
29	2,05	2,76	3,66
30	2,04	2,75	3,65
100	1,96	2,58	3,29

Приложение 2

Таблица 1

**Критические значения выборочного коэффициента корреляции
рангов R_s Спирмена**

z	p		z	p		z	p	
	0,05	0,01		0,05	0,01		0,05	0,01
5	0,94		17	0,48	0,62	29	0,37	0,48
6	0,85		18	0,47	0,60	30	0,36	0,47
7	0,78	0,94	19	0,46	0,58	31	0,36	0,46
8	0,72	0,88	20	0,45	0,57	32	0,36	0,45
9	0,68	0,83	21	0,44	0,56	33	0,34	0,45
10	0,62	0,79	22	0,43	0,54	34	0,34	0,44
11	0,61	0,76	23	0,42	0,53	35	0,33	0,43
12	0,58	0,73	24	0,41	0,52	36	0,33	0,43
13	0,56	0,70	25	0,40	0,51	37	0,33	0,42
14	0,54	0,68	26	0,39	0,50	38	0,32	0,41
15	0,52	0,66	27	0,38	0,49	39	0,32	0,41
16	0,50	0,64	28	0,38	0,48	40	0,31	0,40

Примечание: здесь p – уровень значимости, z – объем выборки.

Если вычисленное значение $R_s < R_{s0,05}$, то корреляция не является статистически значимой. Если эмпирическое значение $R_s \geq R_{s0,01}$, то корреляция является достоверной.

Таблица 2

Критические значения коэффициента корреляции r_{xy} Пирсона

n\P	0,05	0,01	n\P	0,05	0,01
4	0,950	0,990	26	0,388	0,496
5	0,878	0,959	27	0,381	0,487
6	0,811	0,917	28	0,371	0,478
7	0,754	0,874	29	0,367	0,470
8	0,707	0,834	30	0,361	0,463
9	0,666	0,798	35	0,332	0,435
10	0,632	0,765	40	0,310	0,407
11	0,602	0,735	45	0,292	0,384
12	0,576	0,708	50	0,277	0,364
13	0,553	0,684	60	0,253	0,333
14	0,532	0,661	70	0,234	0,308
15	0,514	0,641	80	0,219	0,288
16	0,497	0,623	90	0,206	0,272
17	0,482	0,606	100	0,196	0,258
18	0,468	0,590	125	0,175	0,230
19	0,456	0,575	150	0,160	0,210
20	0,444	0,561	200	0,138	0,182
21	0,433	0,549	250	0,142	0,163
22	0,423	0,537	300	0,113	0,148
23	0,413	0,526	400	0,098	0,128
24	0,404	0,515	500	0,088	0,115
25	0,396	0,505	1000	0,062	0,081

Примечание: корреляция статистически значима, если $r_{xy} \geq r_{xy0,01}$.

Если $r_{xy} < r_{xy0,05}$, то корреляция не является значимой.

Приложение 3

**Примеры применения математико-статистических методов
обработки данных в психологических исследованиях**

Пример 1. Определим, связаны ли между собой индивидуальные показатели готовности к обучению в центре подготовки операторов, полученные до начала обучения у кандидатов на учебу, и их средняя успеваемость в конце периода обучения с помощью коэффициента ранговой корреляции Спирмена.

Коэффициент корреляции R_s Спирмена вычисляется по формуле:

$$R_s = 1 - \frac{\sum_{i=1}^N d_i^2}{N(N^2 - 1)}, \quad (1)$$

где N – количество ранжируемых признаков (показателей, испытуемых); d_i – разность между рангами по двум переменным (например, по качеству деятельности и результатам выполнения теста) для каждого испытуемого.

Для решения этой задачи необходимо проранжировать, во-первых, значения показателей готовности к обучению и, во-вторых, итоговые показатели успеваемости в конце периода обучения. Результаты представлены в табл. 1

Таблица 1

**Ранги показателей готовности к обучению и успеваемости
операторов**

№ кандидата п/п	1	2	3	4	5	6	7	8	9	10	11
Ранги показателей готовности	3	5	6	1	4	11	9	2	8	7	10
Ранги успеваемости	2	7	8	3	4	6	11	1	10	5	9
d_i	1	-2	-2	-2	0	5	-2	1	-2	2	1
d_i^2	1	4	4	4	0	25	4	1	4	4	1

Полученные данные подставляем в формулу (1) и производим расчет:

$$R_s = 1 - 6 \times 52 / (11(11^2 - 1)) = 0,76.$$

Для нахождения уровня значимости обращаемся к табл. 1 (приложение 3), в которой приведены критические значения для коэффициентов ранговой корреляции. Уровни значимости определяем по числу испытуемых n . В нашем случае $n = 11$. Тогда для уровня значимости $p = 0,05$ коэффициент корреляции $R_{кр1} = 0,61$ и для уровня значимости $p = 0,01$ – $R_{кр2} = 0,76$.

Полученный коэффициент корреляции совпал с критическим значением для уровня значимости, равного 0,01. Следовательно, можно утверждать, что показатели готовности и итоговые оценки успеваемости связаны положительной корреляционной зависимостью. Иначе говоря, чем выше показатель готовности, тем успешнее будет учиться кандидат. В терминах статистических гипотез необходимо принять гипотезу о наличии взаимосвязи, т. е. связь между показателями готовности и средней успеваемостью отлична от нуля.

Пример 2. Определим, влияет ли уровень интеллекта операторов на их профессиональные достижения с помощью критерия «хи-квадрат».

С помощью критерия «хи-квадрат» Пирсона $\chi^2_{эмп}$, устанавливающего степень значимости различия распределений признака, полученных при обследовании двух групп лиц, разделенных по частным или обобщенным показателям профессиональной эффективности или другим прогнозируемым показателям, по формуле:

$$\chi^2_{эмп} = \sum_{i=1}^k x_i^2 / f_{in}, \quad (2)$$

где x_i – разность между эмпирическими и «теоретическими» частотами;

k – количество разрядов признака;

f_{mi} – вычисленная или «теоретическая» частота; или

$$\chi^2_{y\ddot{y}i} = M \left\{ \sum_{i=1}^n \sum_{j=1}^m C_j^2 (C_i C_j) - 1 \right\}, \quad (3)$$

где n – число строк многопольной таблицы;

m – число столбцов многопольной таблицы;

M – общее число значений (элементов) в многопольной таблице, вычисляемое по формуле:

$$M = n \times m; \quad (4)$$

C_{ij} – элементы многопольной таблицы;

C_i – суммарные значения по строкам многопольной таблицы;

C_j – суммарные значения по столбцам многопольной таблицы.

При этом статистическая значимость рассчитываемых показателей должна быть не менее 0,05.

Для применения критерия $\chi^2_{\text{эмп}}$ необходимо соблюдать следующие условия:

- измерение может быть проведено в любой шкале;
- выборки должны быть случайными и независимыми;
- желательно, чтобы объем выборки был > 20 . С увеличением объема выборки точность критерия повышается;
- теоретическая частота для каждого выборочного интервала не должна быть меньше 5;
- сумма наблюдений по всем интервалам должна быть равна общему количеству наблюдений;
- таблица критических значений критерия $\chi^2_{\text{эмп}}$ рассчитана для числа степеней свободы ν , которое каждый раз рассчитывают по определенным правилам.

В общем случае число степеней свободы вычисляется по формуле:

$$v = c - 1,$$

(5)

где c – число альтернатив (признаков, значений, элементов) в сравниваемых переменных.

Для таблиц число степеней свободы вычисляют по формуле:

$$v = (n - 1)(m - 1),$$

(6)

где n – число строк, m – число столбцов.

Первый способ решения – по формуле (2).

Например, 90 человек оценили по степени их профессиональных достижений и по уровню интеллекта. При разбиении на уровни (градации признака) по обоим признакам было взято три уровня. Для показателя профессиональных достижений были получены следующие частоты признака: 20 человек с высоким уровнем профессиональных достижений, 40 – со средним и 30 – с низким. Первая группа составляет 22,2 %, вторая – 44,4 % и третья – 33,3 % от всей выборки. При разбиении по уровню интеллекта было взято три равных по численности группы – по 30 человек: с уровнями интеллекта ниже среднего, средним и выше среднего. Каждая группа составляет 33,3 % от всей выборки. Все эмпирические данные (частоты) представлены в табл. 2.

Таблица 2

**Частота распределения испытуемых по уровням
оцениваемых признаков**

IQ	Оценка профессиональных достижений			Всего
	Ниже среднего	Средняя	Выше среднего	
Ниже среднего	20 А (10)	5 S (13,3)	5 C (6,7)	30
Средний	5 D(10)	15 E(13,3)	10 F (6,7)	30
Выше среднего	5 G(10)	20 H(13,3)	5 J (6,7)	30
Итого	30	40	20	90

Для удобства каждая ячейка таблицы обозначена соответствующей латинской буквой: *A*, *S*, *C* и т. д. Табл. 2 устроена следующим образом: в ячейку, обозначенную символом *A*, заносят эмпирические частоты (или число) тех испытуемых, которые одновременно обладают характеристикой: ниже среднего по уровню профессиональных достижений и ниже среднего по интеллекту. Таких испытуемых (эмпирических частот) оказалось 20. В ячейку, обозначаемую символом *S*, заносят эмпирические частоты (или число) тех испытуемых, которые одновременно обладают характеристикой: средние по уровню профессиональных достижений и ниже среднего по интеллекту. Таких испытуемых (эмпирических частот) оказалось 5. В ячейку, обозначенную символом *C*, заносят эмпирические частоты (или число) тех испытуемых, которые одновременно обладают характеристикой: выше среднего по уровню профессиональных достижений и ниже среднего по интеллекту. Таких испытуемых (эмпирических частот) оказалось также 5. Заметим, что $20 + 5 + 5 = 30$, т. е. числу испытуемых, имеющих уровень интеллекта ниже среднего. Подобные «разбиения» были проделаны для каждой ячейки табл. 2. В круглых скобках в каждой ячейке таблицы представлены вычисленные для этой ячейки «теоретические» частоты.

Покажем, как для каждой ячейки табл. 2 найти соответствующую «теоретическую» частоту. Для этого для каждого столбца таблицы подсчитывают так называемые «частости» в процентах:

$$30/90 \cdot 100 \% = 33,3 \%;$$

$$40/90 \cdot 100 \% = 44,4 \%;$$

$$20/90 \cdot 100 \% = 22,2 \%.$$

Полученные величины «частостей» дают возможность подсчитать «теоретические» частоты для каждой ячейки. Они служат основой для подсчета «гипотетических» (а по сути теоретических) частот, т. е. таких, которые при заданном соотношении экспериментальных данных должны были бы быть расположены в соответствующих ячейках табл. 2.

Согласно этому положению «теоретическую» частоту для ячейки *A* подсчитывают следующим образом. 30 человек имеют уровень интеллекта ниже среднего, поэтому 33,3 % от этого числа должны были бы попасть в группу с профессиональными достижениями ниже среднего уровня. Находим эту «гипотетическую» величину: $30 \times 33,3 \% / 100 \% = 9,99 \approx 10$.

Аналогично подсчитывают «теоретические» частоты для ячеек *D, G, S, E, H, C, F* и *J*.

Проверка правильности расчета «теоретических» частот для всех столбцов табл. 2 для ячейки *A* проводят следующим образом: $10 + 10 + 10 = 30$; $13,3 + 13,3 + 13,3 = 39,9 \approx 40$; $6,7 + 6,7 + 6,7 = 20,1 \approx 20$.

Для проверки правильности расчета «теоретических» частот в случае сравнения двух эмпирических наблюдений или для сравнения показателей внутри одной выборки может использоваться следующая формула:

$$f_{cij} = \Sigma f_i \times \Sigma f_j / K, \quad (7)$$

где f_{cij} – теоретическая частота в соответствующей ячейке многопольной таблицы;

Σf_i – сумма эмпирических частот по строке;

Σf_j – сумма эмпирических частот по столбцу;

K – общее количество наблюдений.

Теперь используем формулу (2):

$$\chi^2_{\text{эмп}} = (20-10)^2/10 + (5-13,3)^2/13,3 + (5-6,7)^2/6,7 + (5-10)^2/10 + (15-13,3)^2/13,3 + (10-6,7)^2/6,7 + (5-10)^2/10 + (20-13,3)^2/13,3 + (5-6,7)^2/6,7 = 26,5$$

Число степеней свободы подсчитаем по формуле (6):

$$\nu = (n - 1)(m - 1) = (3 - 1)(3 - 1) = 4.$$

В соответствии с табл. 3 (приложение 1):

$$\chi^2_{\text{кр}1} = 9,49 \text{ для } p \leq 0,05 \text{ и } \chi^2_{\text{кр}2} = 13,28 \text{ для } p \leq 0,01.$$

Полученная эмпирическая величина критерия «хи-квадрат» $\chi^2_{\text{эмп}} = 26,5$ попадает в зону значимости, т. е. $\chi^2_{\text{эмп}} > \chi^2_{\text{кр}2} (p \leq 0,01)$. Иными словами, следует принять гипотезу о том, что уровень интеллекта влияет на успешность профессиональной деятельности.

Второй способ решения – по формуле (3).

Подставив данные табл. 2 в формулу (3), получим:

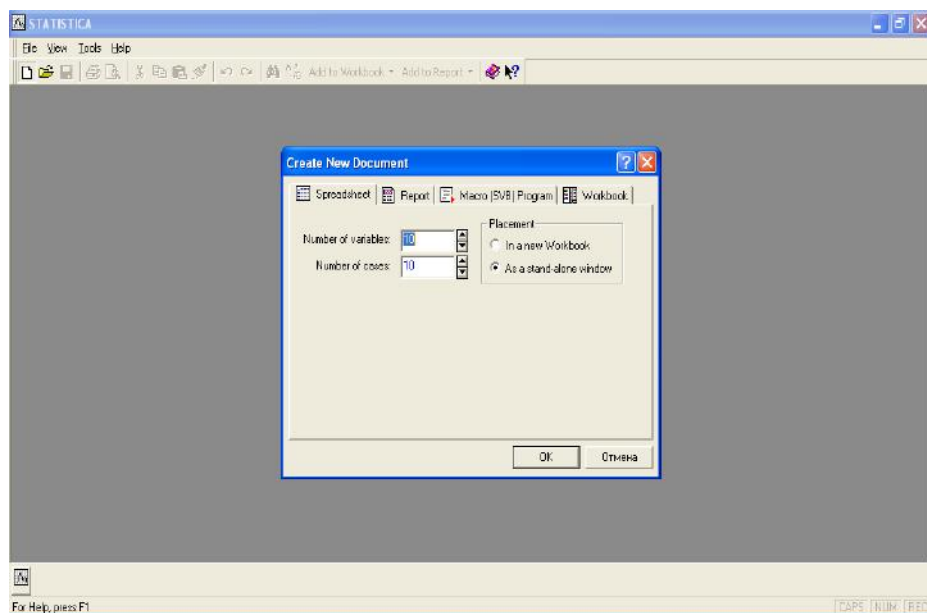
$$\chi^2_{\text{эмп}} = 90(1/30(20^2/30+5^2/40+5^2/20)+1/30(5^2/30+15^2/40+10^2/20)+1/30(5^2/30+20^2/40+5^2/20)-1) = 90(1/2+13/24+1/4-1) = 26,5.$$

Как и следовало ожидать, эмпирическое значение «хи-квадрат» получено то же самое, что и при первом способе решения. Все дальнейшие операции уже проделаны выше при первом способе решения данной задачи. Второй способ существенно проще первого, однако при расчетах по формуле (3) можно легко допустить ошибки. Как первый, так и второй способы расчета эмпирического значения «хи-квадрат» позволяют работать с таблицами практически любой размерности: 3×4 , 4×4 , 5×3 , 5×6 и т. п.

Основные операции с документом программы Statistica 6.0

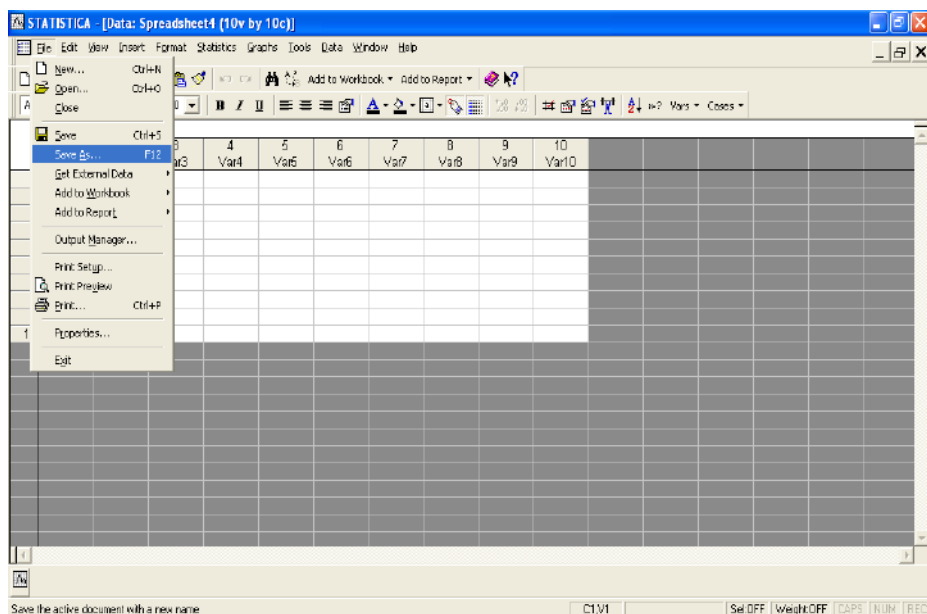
1. Создание нового документа

Нажмите на иконку «New» («Новый документ»). По своему желанию вы можете изменить число переменных (столбцов) и регистров (строчек) создаваемого документа:



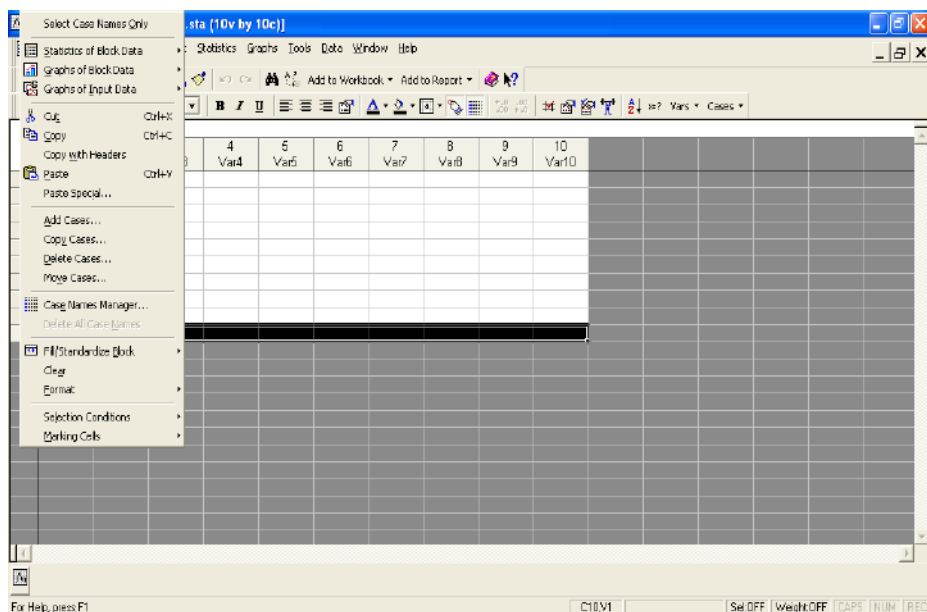
2. Сохранение файла

Нажмите один раз левой кнопкой мышки на функцию «File» («Файл»), выберите операцию «Save As...» («Сохранить как») и назовите свой документ:



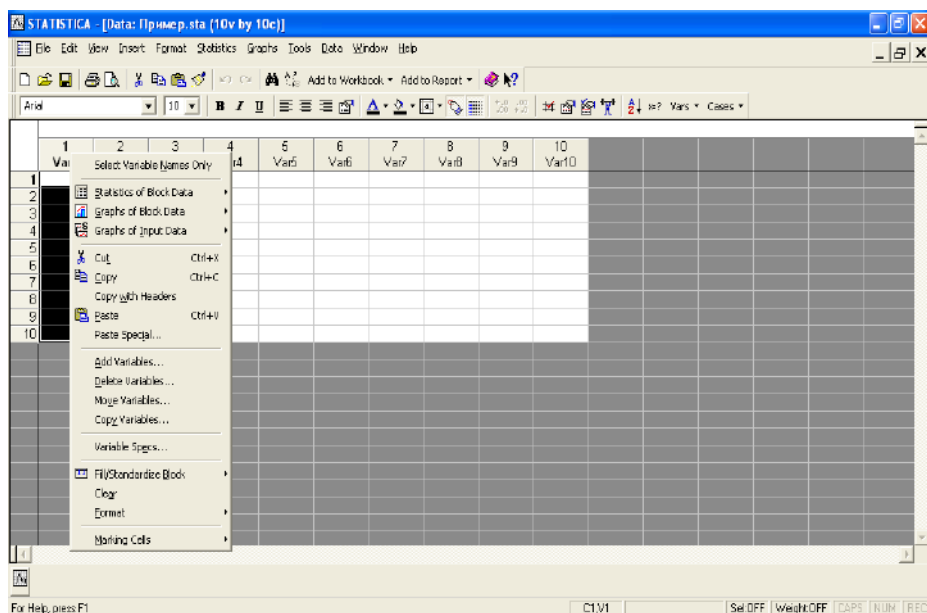
3. Действия с регистром

Нажмите один раз правой кнопкой мышки на регистр таблицы файла (например, нажмите на цифру 10), и перед вами появится окно-меню, в котором представлены различные операции с выбранным регистром. При помощи операции «Add Cases» происходит увеличение строк, «Copy Cases» – копирование строк, «Delete Cases» – удаление строк, «Move Cases» – перемещение строк:



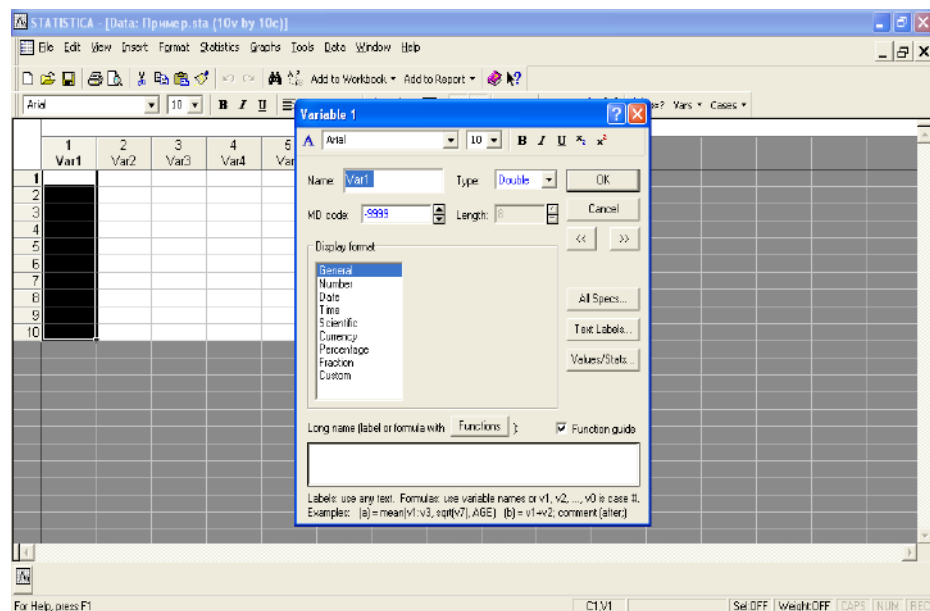
4. Действия с переменной

Нажмите один раз правой кнопкой мышки на название переменной (например, нажмите на Var 1), и перед вами появится окно-меню, в котором представлены различные операции с выбранной переменной:



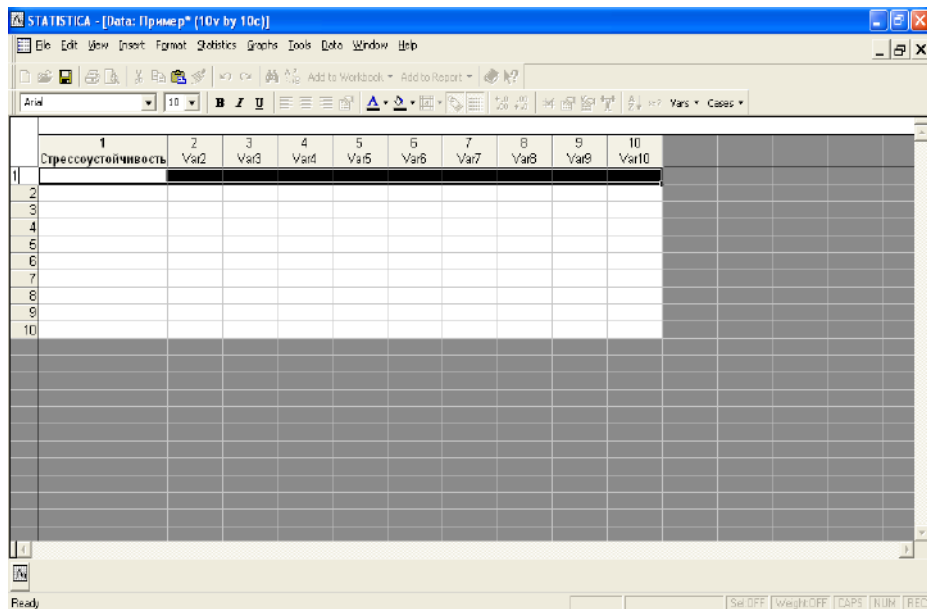
5. Изменение характеристик переменной

Нажмите два раза левой кнопкой мышки на название переменной, и перед вами появится окно-меню, позволяющее изменить различные характеристики выбранного столбца:



6. Изменение названия регистра

Нажмите два раза левой кнопкой мыши на выбранный вами регистр (к примеру, регистр 1):



[illegible]

Учебно-практическое пособие

Богачевский Владимир Александрович,

кандидат психологических наук

Паршутин Игорь Александрович,

кандидат психологических наук, доцент

Сударик Александр Николаевич,

кандидат психологических наук, доцент

ОБРАБОТКА ДАННЫХ ПСИХОЛОГИЧЕСКИХ ИССЛЕДОВАНИЙ С ПОМОЩЬЮ СТАТИСТИЧЕСКИХ ПРОГРАММ



Корректор *Титова В. П.*

Компьютерная верстка *Гриджина Т. А.*

Московский университет МВД России им. В.Я. Кикотя

117997, г. Москва, ул. Академика Волгина, д. 12

Подписано в печать 22.03.2019 г.	Формат 60×84 1/16	Тираж 78 экз.
	Цена договорная	Объем 2,91 уч.-изд. л.
Заказ № 1780		7,2 усл. печ. л.
